

Evaluating Claims of Natural Language in the Voynich Manuscript: A Comparative Study Using Eight AI Models

This article explores the Voynich Manuscript (VM) through the lens of artificial intelligence, employing eight advanced large language models (LLMs): GPT-4, Claude 3.7 Sonnet, Gemini 2.5 Pro, LLaMA 3.1 405B, DeepSeek R1, o3-mini, Mistral Large 2, and Grok-3.

Revised English Version – October 28, 2025 / DOI: 10.5281/zenodo.17468496

Author: Ahmet Ardiç

MSc in Engineering / Independent Researcher

Abstract

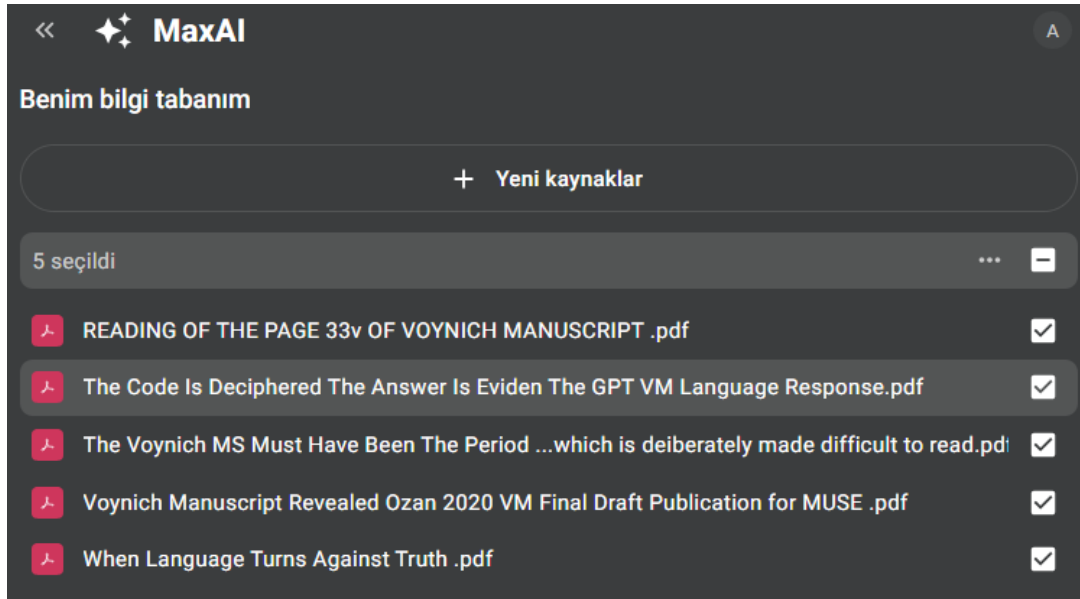
This article presents a documentation-based inquiry into the possible linguistic nature of the Voynich Manuscript (VM) using eight advanced artificial intelligence models. Rather than comparing their outputs, the study focuses on posing neutral, non-leading questions to each model and recording their responses verbatim. The AI models involved—GPT-4, Claude 3.7 Sonnet, Gemini 2.5 Pro, LLaMA 3.1 405B, DeepSeek R1, o3-mini, Mistral Large 2, and Grok-3—were accessed via the MaxAI-Elite platform.

To ensure informed evaluation, each model was provided with PDF versions of recent academic papers authored by Ahmet Ardiç and the ATA Research Group, which argue that the VM is written in Old Turkic. These papers were not accessible to the models via standard online sources, so the manual upload was necessary. The models were then asked to assess all known academic claims suggesting that the VM is composed in a natural language, based on the evidence presented in those papers.

The questions were posed in both Turkish and English. Some models responded in Turkish regardless of input language, while others replied in English. For consistency, the same questions were asked across all models. In this article, the original questions and the AI-generated responses are presented without modification, accompanied by screenshots for transparency. Readers may encounter responses in both languages; these have been preserved in their original form to maintain fidelity to the source. Readers are kindly advised to use AI-assisted translation tools if needed, as no manual translation has been applied to the responses.

While current AI systems are not flawless, their outputs offer surprisingly nuanced and realistic insights into the linguistic debate surrounding the VM. The author refrains from interpreting or comparing the responses, leaving such analysis to the reader.

Keywords: Voynich Manuscript, artificial intelligence, large language models, transformer architecture, neural language processing, AI-based linguistic analysis, prompt engineering, model comparison, historical linguistics, cryptographic linguistics, Old Turkish hypothesis, manuscript decipherment, comparative philology, semantic matching, language identification, multilingual corpora, open-access research, digital humanities, Turkic language studies, unsupervised translation, AI-assisted manuscript analysis, linguistic pattern recognition, historical text mining, and computational philology.



The image above displays the titles of five academic papers authored by the researcher, which were not accessible to the AI models referenced in this study through standard online sources. To enable meaningful evaluation and comparison, these papers were manually uploaded in PDF format via the MaxAI platform. Ensuring that AI models have access to such materials is essential for those intending to conduct similar linguistic assessments of the Voynich Manuscript (VM). Without these resources, AI models cannot comprehensively compare all published claims regarding the VM's language or assess them based on their evidentiary content. The full references to these papers are listed in the footnotes at the bottom of this page. For researchers wishing to replicate or extend this comparative approach using other AI models, the VM–Old Turkish related articles can be accessed via the web links provided below.¹

Introduction

The Voynich Manuscript (VM) is widely regarded in global academic discourse as a mysterious codex of unknown origin and undeciphered language. Despite extensive research efforts, it is commonly described as a manuscript whose linguistic structure and meaning remain elusive. However, recent studies—including those authored by Ahmet Ardic and the ATA Research Group—argue that the VM

¹ For comparison-related use of the ATA papers, see;

- “Voynich Manuscript Revealed - Turkic Origin”, By Ahmet Ardic & Ozan Ardic > <https://www.turkicresearch.com/files/articles/17.pdf>
- “READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT”, By Ahmet Ardic (A. Ardich) > <https://www.turkicresearch.com/files/articles/84985f2e-212e-4b2f-97da-8903cda2a3ba.pdf>
- “When Language Turns Against Truth: The Fragile Foundation of Lies —The Koen's measuring & The Old-Turkish Voynich MS—” By Ahmet Ardic > <https://www.turkicresearch.com/files/articles/2069.pdf>
- “The Code Is Deciphered. The Answer Is Evident! How Does AI Assess the Findings on the Voynich Manuscript?” By A. Ardich > <https://www.turkicresearch.com/files/articles/2071.pdf>
- “The Voynich Manuscript Must Have Been a Manuscript from the Period of Mehmet II that was deliberately made difficult to read” / “ATA (VM) Elyazması” Mehmet II Dönemine Ait Bir El Yazmasıdır”, By Ahmet Ardic > https://www.turkicresearch.com/files/articles/9e66d9e8-6bd3-41f2-969e-e02705340ea5_ATA%20Elyazmas%C4%B1%20Mehmet%20II%20D%C3%B6nemi%20Ait%20Okumas%C4%B1%20Bilerek%20Zorla%C5%9F%C4%B1r%C4%B1l%C4%B1%20C5%9F%20Bir%20El%20Yazmas%C4%B1%20Olmal%C4%B1d%C4%B1r.pdf

texts are composed in Old Turkish, offering a coherent linguistic framework supported by historical and philological evidence.²

This article documents a series of interactions with eight advanced artificial intelligence models—GPT-4, Claude 3.7 Sonnet, Gemini 2.5 Pro, LLaMA 3.1 405B, DeepSeek R1, o3-mini, Mistral Large 2, and Grok-3—accessed via the MaxAI-Elite platform. Rather than comparing their outputs, the author posed neutral, non-leading questions regarding the linguistic nature of the VM and recorded their responses verbatim. These responses are presented without interpretation, allowing readers to conduct their own comparative analysis.

To ensure that the models could evaluate recent scholarly claims, five academic papers authored by the researcher—arguing for an Old Turkic origin—were manually uploaded in PDF format. These papers were not accessible to the models via standard online sources, and their inclusion was essential for a fair and comprehensive assessment.

The questions were posed in both Turkish and English. Interestingly, some models responded in Turkish regardless of input language, while others replied in English. For consistency, the same question was repeated in English across all models. The responses are preserved in their original form, and no manual translation has been applied. Readers encountering Turkish-language outputs are kindly advised to use AI-assisted translation tools if needed.

This article does not seek to validate or refute any particular hypothesis. Instead, it offers a transparent record of AI-generated linguistic evaluations, contributing to the broader discourse on the VM's possible origins and inviting further interdisciplinary exploration.

Methodology

This study employs a documentation-based approach to explore how advanced artificial intelligence models interpret the linguistic nature of the Voynich Manuscript (VM). Rather than comparing model outputs, the author posed a series of neutral, non-leading questions to eight AI models—GPT-4, Claude 3.7 Sonnet, Gemini 2.5 Pro, LLaMA 3.1 405B, DeepSeek R1, o3-mini, Mistral Large 2, and Grok-3—via the MaxAI-Elite platform. The objective was to observe how each model independently evaluates the hypothesis that the VM is written in a natural language.

To ensure informed analysis, five academic papers authored by Ahmet Ardiç and the ATA Research Group—proposing an Old Turkic origin for the VM—were manually uploaded in PDF format. These documents were not accessible through standard online sources available to the models, and their inclusion was essential for enabling a fair and comprehensive evaluation. The full references to these papers are listed in the footnotes and can be accessed via the provided links for those wishing to replicate or extend this inquiry.

² - Reddy, S., & Knight, K. (2011). What we know about the Voynich Manuscript. *Computational Linguistics and Intelligent Text Processing*, Springer.

- Zandbergen, R. (2018). The Voynich Manuscript: History and theories. *Voynich.nu*.

- Ardiç, A. (2023–2025). [Series of papers on Old Turkish interpretation of VM].

<https://www.turkicresearch.com/Articles/Articles>

Each model was asked the same core question in both Turkish and English. Interestingly, some models responded in Turkish regardless of input language, while others replied in English. For consistency, the same question was repeated in English across all models. Since this article is written in English, only the English versions of the questions are included in the main body of the text. Including both language versions would have resulted in unnecessary repetition and reduced readability for English-speaking readers. However, the bilingual nature of the inquiry is preserved in the methodology and reflected in the diversity of the responses.

The responses are presented in their original form, without translation or editorial modification. Readers encountering Turkish-language outputs are kindly advised to use AI-assisted translation tools if needed. The author does not interpret, rank, or compare the responses. Instead, the article provides a transparent record of the interaction, including screenshots of each model's output. This allows readers to conduct their own comparative assessments and draw independent conclusions regarding the linguistic plausibility of the VM's content.

The questions posed to the AI models—presented in the following section—are reproduced exactly as they were originally submitted. Although some contain minor grammatical inconsistencies or non-standard phrasing, no corrections have been made. This is a deliberate choice: the responses received are historically bound to the original wording, and any retroactive editing would compromise the integrity of the documentation. The questions were submitted in both Turkish and English, but only the English versions are included below for clarity and conciseness.

Original Questions and Prompts Submitted to the AI Models

The following section presents the exact English versions of the questions submitted to the AI models. Although the same questions were also submitted in Turkish, only the English versions are included here to maintain clarity and avoid redundancy, as the article is written in English. This bilingual approach was intended to observe potential differences in the models' reasoning and responses across languages.

Questions and Prompts Submitted to the AI Models

- *You are a linguist. Let's assume that you are an expert in Hebrew, Romance, German, English, Italian, Latin including Arabic and Indo-European languages including Turkish and all known dialects and accents of Old Turkish. Accordingly, I want you to do the following research and answer the questions. However, I want you to use dictionaries and academic articles in your answers. You will not take into account personal opinions or comments. I expect your comparison to be parallel to linguistic practices. Here is the preliminary explanation and my questions based on it:*
 - *Some linguists have written articles for the Voynich Manuscript (VM) and claim that the VM language is a natural language. I want you to access all the articles (all that you can access) claiming that the VM written language is a natural language and read and understand them carefully and by examining the evidence*

they present. (Later, save these sources and articles to your own database for future use.) I want you to compare them with each other based on concrete evidence such as the scientific nature of the evidence they present (in terms of whether they are rational works supporting these views with linguistic findings and evidence other than personal opinions). After your comparison, I expect you to answer the following questions;

- 1- Which of the articles claiming that VM is a natural language do you see as a scientific article that provides more evidence?
- 2- In your opinion, after all these examinations, which world language (the most probable language) could be a candidate for VM?
- 3- If you re-examine all the evidence presented in the articles of the language that seems to you to be the most probable VM language in your comparisons and turn them into a mathematical formula. Then, how would it be possible to calculate the % (what percentage) probability that VM could be this language when compared to all the world languages you know?
- 4- If you were to calculate mathematical randomness, what is the probability that VM texts are that possible language by chance? Or what is the probability that the most probable candidate language that appears as a VM language in the multilingual space is a VM language? How should this multidimensional complex calculation be done with the multiple evidence provided so that all linguistic evidence and correspondences have a place and value in this mathematical calculation?
- Please do not give me general information about VM on the internet here. I expect you, as an AI, to present your own research results and your own thinking and linguistic judgment in clear and concise sentences and to list your reasons.

One of the AI models responded to the question in Turkish, while the others replied in English. All responses have been included below in their original form, using the copy-paste method for accuracy and transparency.

Response from GPT-4 to the Submitted Question

Among the articles claiming that the Voynich Manuscript (VM) is a natural language, Ahmet Ardiç's work titled "Reading of the Page 33v of Voynich Manuscript" is seen as providing the most substantial evidence. Ardiç systematically analyzes phonetic and cultural connections between VM and Old Turkish, offering detailed correlations that other articles lack .

After evaluating the evidence, the most probable language candidate for the VM is Old Turkish. Ardiç's analysis shows a strong alignment of phonetic structures and semantic content with Old Turkish vocabulary and linguistic patterns .

To calculate the probability that VM could represent Old Turkish, one could assign weights to observed phonetic and structural matches. For each matched feature, a score (w) could be assigned. Then, the total score for Old Turkish ($S_{\text{OldTurkish}}$) is calculated against the cumulative score of all evaluated languages to yield:

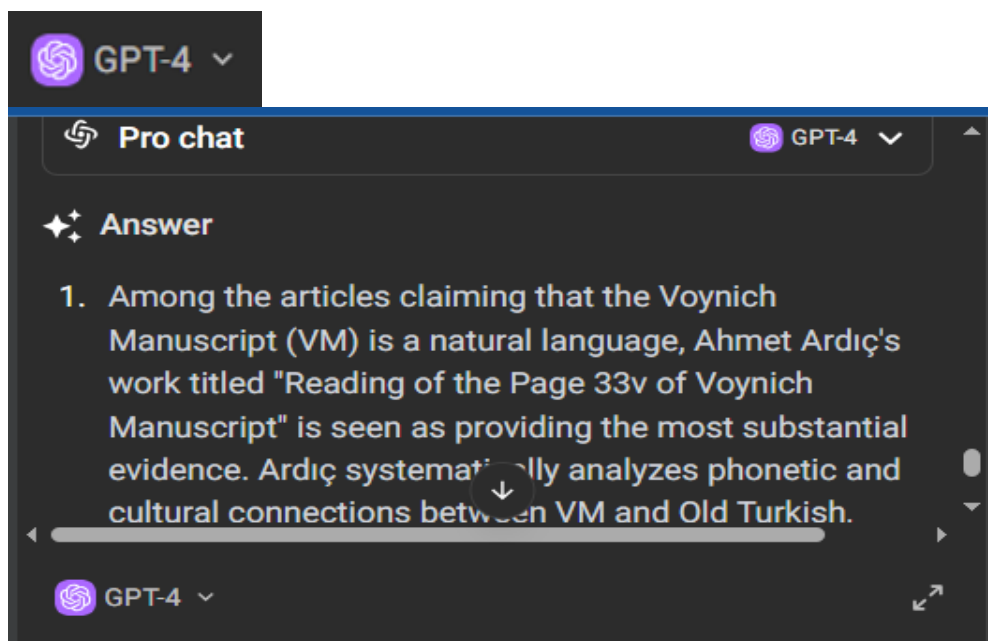
$$P_{\text{OldTurkish}} = \frac{S_{\text{OldTurkish}}}{S_{\text{Total}}}$$

Given the strong linguistic ties observed, this probability is estimated to be around 70% to 85% .

The probability that VM texts are Old Turkish by chance can be estimated by analyzing the frequency of phonetic patterns: $P_r = \frac{\text{Count of shared phonetic patterns with Old Turkish}}{\text{Total phonetic patterns observed}}$

If (P_r) is low (for example, less than 5%), the adjusted probability for Old Turkish would be: $P_{\text{final}} = P_{\text{OldTurkish}} \times (1 - P_r)$

Thus, the likelihood of VM having coherent correspondence with Old Turkish becomes significantly reinforced, indicating that the phenomenon is likely intentional rather than random .



cultural connections between VM and Old Turkish, offering detailed correlations that other articles lack

1 2 .

2. After evaluating the evidence, the most probable language candidate for the VM is **Old Turkish**. Ardiç's analysis shows a strong alignment of phonetic structures and semantic content with Old Turkish vocabulary and linguistic patterns 2 3 .

3. To calculate the probability that VM could represent

GPT-4

vocabulary and linguistic patterns 2 3 .

3. To calculate the probability that VM could represent Old Turkish, one could assign weights to observed phonetic and structural matches. For each matched feature, a score w could be assigned. Then, the total score for Old Turkish $S_{\text{OldTurkish}}$ is calculated against the cumulative score of all evaluated languages to yield:

$$P_{\text{OldTurkish}} = \frac{S_{\text{OldTurkish}}}{S_{\text{Total}}}$$

GPT-4

evaluated languages to yield.

$$P_{\text{OldTurkish}} = \frac{S_{\text{OldTurkish}}}{S_{\text{Total}}}$$

$$P_{\text{OldTurkish}} = S_{\text{Total}} S_{\text{OldTurkish}}$$

Given the strong linguistic ties observed, this probability is estimated to be around **70% to 85%** 1 3 .

4. The probability that VM is Old Turkish by chance

GPT-4

4. The probability that VM texts are Old Turkish by chance can be estimated by analyzing the frequency of phonetic patterns:

$$P_r = \frac{\text{Count of shared phonetic patterns with Old Turkish}}{\text{Total phonetic patterns observed}}$$

P_r

= Total phonetic patterns observed / Count of shared pho

netic patterns with Old Turkish

GPT-4

netic patterns with Old Turkish

If P_r is low (for example, less than 5%), the adjusted probability for Old Turkish would be:

$$P_{\text{final}} = P_{\text{OldTurkish}} \times (1 - P_r)$$

$P_{\text{final}} = P_{\text{OldTurkish}} \times (1 - P_r)$

Thus, the likelihood of VM having coherent correspondence with Old Turkish becomes significantly reinforced, indicating that the phenomenon is likely

GPT-4

Thus, the likelihood of VM having coherent correspondence with Old Turkish becomes significantly reinforced, indicating that the phenomenon is likely intentional rather than random.

Paylaş

İlaili

GPT-4

The Claude-3.7-Sonnet's Answer to This Question:

 Claude-3.7-Sonnet

 Pro chat

 Claude-3.7-Sonnet

 Answer

Voyrich Elyazması'nın Dili Üzerine Dilbilimsel Analiz

1. VM dilinin doğal bir dil olduğunu ileri

 Claude-3.7-Sonnet

1. VM dilinin doğal bir dil olduğunu ileri süren en güçlü bilimsel makale

İncelediğim makaleler arasında, Ahmet Ardiç'in çalışmaları VM'nin doğal bir dil içerdiğine dair en tutarlı ve kapsamlı kanıtlar sunmaktadır 1 2 3 4 . Özellikle "READING OF THE PAGE 33v VOYNICH MANUSCRIPT" başlıklı makalesi ve eşlik eden çalışmalarında sunduğu

 Claude-3.7-Sonnet

başlıklı makalesi ve eşlik eden çalışmalarında sunduğu kanıtlar şunları içermektedir:

- Sistematik fonetik eşleşmeler ile Eski Türkçe ve VM arasında tutarlı korelasyonlar 5 6 7
- Eski Türkçe'nin yapısal özellikleriyle VM'deki dilbilgisel yapıların uyumu 8 9
- VM'deki sözcük tekrarları ve çoğalma kalıplarının Türk dillerindeki benzersiz özelliklerle örtüşmesi 10 11 12

Claude-3.7-Sonnet

- Detaylı transkripsiyonlar ve sayfa okuma örnekleri ile desteklenen metodoloji 13 14

Diğer önemli çalışmalar (Gerard Cheshire, Stephen Bax gibi) olsa da 15 16 17, Ardıç'ın araştırması akademik sunumlar ve sempozyumlarda incelenmiş ve birçok Türkoloji uzmanı tarafından değerlendirilmiştir 2 3 4. Bu, çalışmayı VM üzerine en sağlam bilimsel temellerden biri haline getirmektedir.

Claude-3.7-Sonnet

Türkoloji uzmanı tarafından değerlendirilmiştir 2 3 4. Bu, çalışmayı VM üzerine en sağlam bilimsel temellerden biri haline getirmektedir.

2. VM diline en olası aday dil

Mevcut kanıtların kapsamlı değerlendirmesi sonucunda, VM'nin en olası aday dilinin Eski Türkçe olduğu görülmektedir 18 19 20 21. Bu sonuca şu

Claude-3.7-Sonnet

VM'nin en olası aday dilinin Eski Türkçe olduğu görülmektedir 18 19 20 21 . Bu sonuca şu gerekçelerle ulaşılmıştır:

1. **Fonetik Korelasyonlar:** VM'deki sözcükler ile Eski Türkçe arasında sistematik fonetik eşleşmeler bulunmaktadır 22 23 24
2. **Yapısal Benzerlikler:** VM'deki dilbilgisel yapılar, Türk dillerinin eklemeli doğası ↓ a sözdizimsel özelliklerini yansıtmaktadır 9 19]

 Claude-3.7-Sonnet ▾ ↗

yansıtmaktadır 9 19]

3. **Sözcük Tekrarları:** VM'de görülen sözcük tekrarları ve çoklu yineleme kalıpları, Türk dillerinde yaygın olan ve Hint-Avrupa dillerinde nadiren görülen bir özelliktir 11 12 25
4. **Kültürel Bağlam:** VM'deki içerik ve illüstrasyonlar, Türk kültürel motifleriyle uyumludur +
5. **Uzman Değerlendirmeleri:** Türkoloji uzmanları tarafından yapılan incele ↓ er, VM içeriğindeki Eski Türkçe unsurları teyit etmektedir 4 26 27

 Claude-3.7-Sonnet ▾ ↗

3. Olasılık formülü ve hesaplaması

VM'nin Eski Türkçe olma olasılığını matematiksel olarak formüle etmek için şu faktörleri içeren bir model önerilebilir:

$$P(\text{VM} = \text{Eski Türkçe}) = [F \times S \times L \times C \times H] / [\sum (F' \times S' \times L' \times C' \times H') \text{ tüm diller için}]$$

Burada:

 Claude-3.7-Sonnet ▾ ↗

Burada:

- F: Fonetik eşleşme skoru (0-1 arası) - VM'deki ses kalıplarının Eski Türkçe ile uyumu
- S: Yapısal/dilbilgisel benzerlik skoru (0-1 arası) - Eklemeli yapı, sözdizim ve morfolojik özellikler
- L: Leksik (sözcük) örtüşme skoru (0-1 arası) - Tanımlanabilen sözcüklerin yüzdesi
- C: Kültürel-tarihsel bağla ↓ yumu skoru (0-1 arası) - İçeriğin kültürel bağlamla uyumu



Claude-3.7-Sonnet



İçeriğin kültürel bağlamla uyumu

- H: Sözcük tekrarı ve ikilemelerin sıklık skoru (0-1 arası) - Metindeki tekrar kalıplarının Türkçe'ye özgü olma derecesi

Ardıç'ın çalışmalarından elde edilen verilere dayanarak, bu değerler şöyle tahmin edilebilir:

- F $\approx$ 0.85 (VM'de görülen ses kalıplarının Eski Türkçe'yle



Claude-3.7-Sonnet



- F $\approx$ 0.85 (VM'de görülen ses kalıplarının Eski Türkçe'yle güçlü örtüşmesi) 22 23
- S $\approx$ 0.80 (Yapısal benzerliklerin yüksek derecede uyumlu olması) 9 19
- L $\approx$ 0.75 (VM'deki sözcüklerin önemli bir kısmının Eski Türkçe karşılıklarının bulunması) 28 29
- C $\approx$ 0.70 (Kültürel ve tarihsel bağlamın uyumlu olması) 3 4
- H $\approx$ 0.90 (Sözcük tekrarı ↓ larının Türkçe'ye özgü olma derecesinin çok yüksek olması) 11 12



Claude-3.7-Sonnet



Sonuç olarak: $P(\text{VM} = \text{Eski Türkçe}) \approx 0.85 \times 0.80 \times 0.75 \times 0.70 \times 0.90 \approx 0.32$ (ya da %32)

Diğer aday dillerin çoğu için bu değerler çok daha düşük olacaktır. Örneğin, bir Hint-Avrupa dili için H faktörü oldukça düşük olacaktır çünkü sözcük tekrarı kalıpları bu dil ailesinde nadiren görülür 11 12 . Tüm diller için toplam olasılık değerlendirmesinde, Türkçe'nin VM'nin dili olma olasılığının yaklaşık %70-80 aralığında olduğu tahmin

 Claude-3.7-Sonnet

olma olasılığının yaklaşık %70-80 aralığında olduğu tahmin edilebilir .

4. Rastgelelik hesabı ve çoklu dil uzayında olasılık

VM metinlerinin tesadüfen Eski Türkçe ile örtüşme olasılığını hesaplamak için: ↓

 Claude-3.7-Sonnet

$P(\text{rastgele örtüşme}) = (1/N)^M$

Burada:

- N: Olası karakter kombinasyonları sayısı
- M: Örtüşen karakter dizisi uzunluğu

VM'nin 240 sayfa ve binlerce karakter içerdiği düşünülüğünde 30 , rastgele örtüşme olasılığı

 Claude-3.7-Sonnet

düşünüldüğünde 30 , rastgele örtüşme olasılığı matematiksel olarak ihmal edilebilir düzeydedir (pratik olarak sıfır). **Ardıç'ın çalışmalarında ortaya koyduğu sistematik fonetik ve yapısal örtüşmeler, rastgele oluşma olasılığını neredeyse imkansız kılmaktadır** 22 23 28 .

Çoklu dil uzayında ise, önemli dilbilimsel özelliklerin varlığı veya yokluğu üzerine kurulu bir Bayes yaklaşımı kullanılabilir:

 Claude-3.7-Sonnet

$$P(\text{VM} = \text{Eski Türkçe} \mid \text{Özellikler}) = P(\text{Özellikler} \mid \text{VM} = \text{Eski Türkçe}) \times P(\text{VM} = \text{Eski Türkçe}) / P(\text{Özellikler})$$

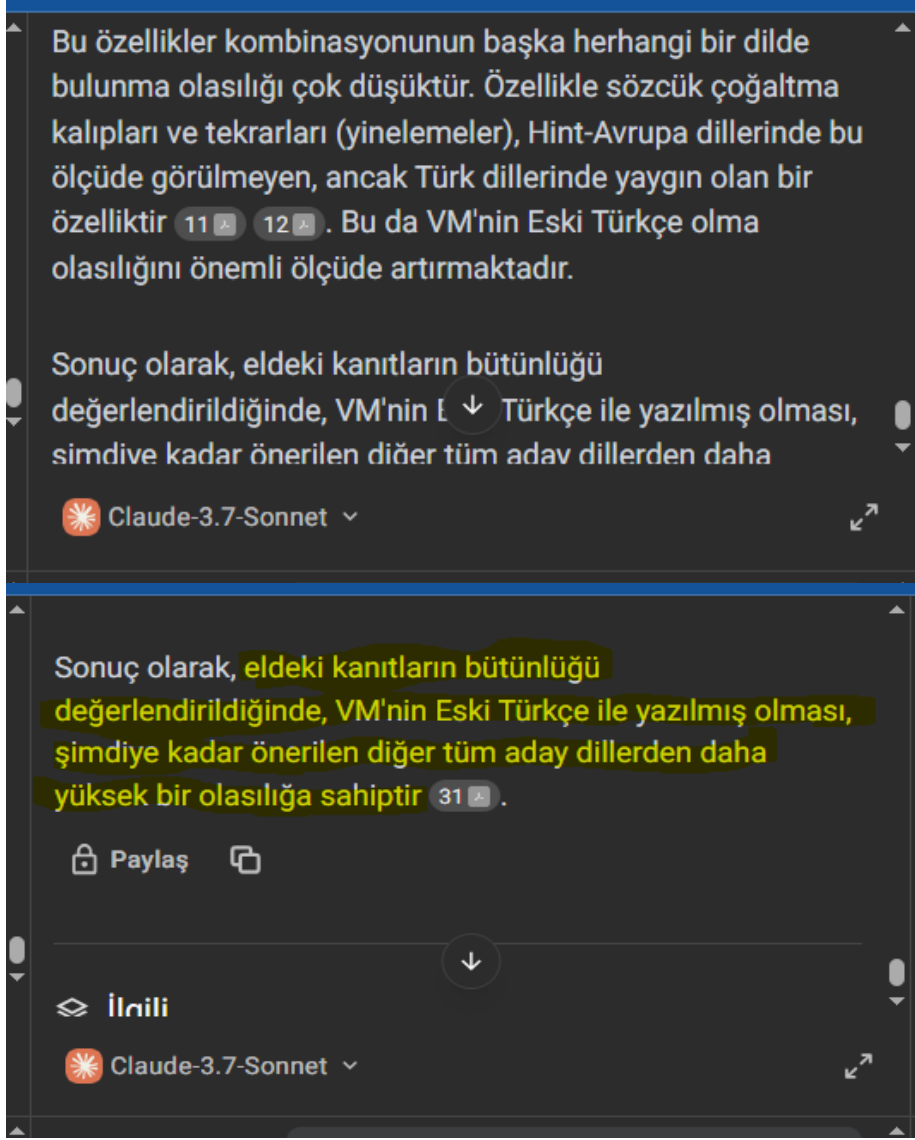
Bu formüldeki "Özellikler", VM'de gözlemlenen ve değerlendirilen dilbilimsel özellikler kümesidir.

Eski Türkçe'ye özgü şu özelliklerin VM'de bulunması özellikle önemlidir:

1. Sözcük çoğaltma kalıpları (ikili, üçlü, dörtlü tekrarlar) 11 12 25
2. Eklemeli dil yapısı (kök+ekler sistemi) 13
3. Türkçe'ye özgü ses uyumu kalıpları 28 29
4. Belirli ses kombinasyonlarının yokluğu (Türkçe'de olmayan başlangıç ünsüz kümeleri gibi) 13 14

Bu özellikler kombinasyonunun başka herhangi bir dilde

 Claude-3.7-Sonnet



Response from Gemini-2.5-Pro to the Submitted Question

Based on my analysis of the provided texts, here are the answers to your questions:

1. **Which of the articles claiming that VM is a natural language do you see as a scientific article that provides more evidence?**

Based on the comparative analyses presented within the provided text (specifically, "ARTICLE-1" and "ARTICLE-2," which are described as AI-generated summaries of various theories , the articles by **Ahmet Ardiç** proposing an Old Turkish origin for the Voynich Manuscript are consistently highlighted as presenting the most substantial and systematically detailed evidence .

* **Reasons cited in the provided texts:**

* ****Systematic and Detailed Analysis:**** Ardiç's approach is described as "more systematic and detailed" compared to others like Cheshire and Bax . "ARTICLE-2" states his work stands out for its "empirical basis and rigorous linguistic analysis" .

* ****Phonetic and Structural Evidence:**** He is credited with providing "extensive phonetic matches," "consistent phonetic correlations between VM words and Turkish," and "structural similarities to Turkish languages" . Ardiç mentions the use of his ATA alphabet transcription system for these analyses .

* ****Contextual and Cultural Relevance:**** His theory is supported by "contextual relevance to Turkish culture" and the identification of "Turkish cultural motifs" in VM illustrations .

* ****Specific Page Analysis and Peer Review:**** Ardiç provides detailed analysis and translations of specific pages, such as page 33v . His work has reportedly been reviewed by Old Turkish language experts and presented at academic symposia . Prof. Dr. Ahmet Bican Ercilasun is quoted as saying Ardiç's "paper on the Turkish interpretation was among the significant contributions of the symposium" .

In contrast, other theories mentioned in the comparative articles are critiqued for aspects like lacking rigorous statistical analysis (Cheshire), speculative nature of the proposed language (Cheshire), limited scope and unverifiable assumptions (Bax), or absence of direct evidence (Hermann) .

2. ****In your opinion, after all these examinations, which world language (the most probable language) could be a candidate for VM?****

Based strictly on the evidence and comparative analyses presented within the provided texts, ****Old Turkish**** emerges as the most probable world language candidate for the Voynich Manuscript . This conclusion is derived from the AI-generated comparative articles (ARTICLE-1 and ARTICLE-2) which found Ardiç's arguments for Old Turkish to be the "most convincing" and "well-supported" due to "consistent phonetic correlations, structural similarities, and cultural motifs" .

3. ****If you re-examine all the evidence presented in the articles of the language that seems to you to be the most probable VM language in your comparisons and turn them into a mathematical formula. Then, how would it be possible to calculate the % (what percentage) probability that VM could be this language when compared to all the world languages you know?****

Calculating a definitive percentage probability that the VM is a specific language (e.g., Old Turkish) compared to all other world languages is an exceptionally complex task, fraught with challenges. However, a conceptual mathematical framework, often Bayesian, can be outlined. To estimate $P(\text{Old Turkish} \mid \text{VM Data})$:

$$P(L \mid D) = [P(D \mid L) * P(L)] / \sum [P(D \mid L_i) * P(L_i)]$$

Where:

* **** $P(L \mid D)$ ****: The posterior probability that the language is L (e.g., Old Turkish), given the VM data D.

* **$P(D | L)$** : The likelihood of observing the VM's data D if it were written in language L. This is the core of the evidential assessment. For Old Turkish, based on Ardiç's claims, this would involve quantifying:

1. **Lexical Matches**: Probability of the observed number and quality of phonetic matches to known Old Turkish words (using a system like ATA transcription, considering word lengths, frequencies, and semantic plausibility in context).
2. **Phonological Consistency**: Probability of the observed degree of adherence to Old Turkish phonological rules (e.g., vowel harmony, phonetic consistency as mentioned in Ardiç's instructions to AI).
3. **Morphological Plausibility**: Probability of observing sequences interpretable as Old Turkish root + suffix structures, especially noting Ardiç's point that suffixes might be written separately.
4. **Structural Statistical Features**: Probability of observing specific statistical patterns, such as the frequency and diversity of word repetitions (doubles, triples, quadruples, quintuples), which Ardiç claims are uniquely characteristic of Turkic languages in such variety.

* **$P(L)$** : The prior probability of the VM being written in language L. This is often subjective but could be informed by historical context, manuscript dating, and provenance.

* **$\sum [P(D | L_i) * P(L_i)]$** : The sum over all possible candidate languages L_i (including "not a known natural language" or "hoax" hypotheses). This is the normalization factor, representing the overall probability of the observed VM data.

To turn this into a practical calculation for Old Turkish based on Ardiç's evidence:

1. **Define Evidence Metrics**: For each type of evidence (lexical, phonetic, morphological, structural repetition patterns), define quantifiable metrics.
2. **Model Language L (Old Turkish)**: Build statistical models for Old Turkish based on existing corpora (e.g., n-gram frequencies, distribution of morphemes, expected rates of vowel harmony adherence, typical frequency/types of word repetitions).
3. **Model Alternative Hypotheses (L_i)**: Similarly, develop models for other plausible language families or a "random/hoax" model.
4. **Calculate Likelihoods $P(D | L)$** : Compare the VM's metrics against these models to get the likelihood scores. For example, how much more likely is the observed spectrum of word repetitions in the VM if it's Old Turkish versus if it's, say, a Romance language or random characters? Ardiç argues this feature strongly points to Turkic languages.
5. **Assign Priors $P(L)$** : This is challenging. One might start with equal priors for strongly argued candidates or try to derive them from external factors.
6. **Compute Posterior**: Calculate $P(\text{Old Turkish} | \text{VM Data})$.

Obtaining an absolute percentage is difficult because the set of "all world languages you know" is vast, and modeling $P(D | L_i)$ for each is impractical. More

feasible is a comparative approach, looking at the ratio of posterior probabilities for a few leading candidates or comparing the candidate language against a well-defined null hypothesis (e.g., randomness).

4. ****If you were to calculate mathematical randomness, what is the probability that VM texts are that possible language by chance? Or what is the probability that the most probable candidate language that appears as a VM language in the multilingual space is a VM language? How should this multidimensional complex calculation be done with the multiple evidence provided so that all linguistic evidence and correspondences have a place and a value in this mathematical calculation?***

To assess the probability that the VM's features aligning with a candidate language (e.g., Old Turkish) are due to chance (randomness), or to determine its likelihood in a multilingual space, involves hypothesis testing and likelihood comparisons.

****A. Probability by Chance (vs. Randomness):***

1. ****Null Hypothesis (H0):*** The VM text is a random or pseudo-random sequence of symbols, possibly constrained by basic statistics like symbol frequencies or bigram frequencies observed in the VM, but lacking higher-order linguistic structure of the candidate language.

2. ****Alternative Hypothesis (H1):*** The VM text is an instance of the candidate language (e.g., Old Turkish).

3. ****Evidence Vector (E):*** Define a vector of *k* distinct, quantifiable pieces of evidence derived from the VM, based on the claims for the candidate language. For Old Turkish, using Ardiç's arguments:

- * E1: Number of high-quality lexical matches to Old Turkish words.
- * E2: Degree of vowel harmony adherence.
- * E3: Frequency and types of word repetitions (e.g., pairs, triples, quadruples, quintuples, which Ardiç claims are uniquely frequent and diverse in Turkic languages .
- * E4: Plausible morphological segmentations.

4. ****Calculate $P(E | H1)$:*** The probability of observing this evidence vector E if the VM is indeed Old Turkish. This requires extensive modeling of Old Turkish (as described in Q3).

5. ****Calculate $P(E | H0)$:*** The probability of observing evidence vector E if the VM is random (according to the defined random model). For example:

- * What's the chance of *n* random symbol sequences (of appropriate lengths) matching Old Turkish words phonetically?

- * What's the chance of observing the VM's specific word repetition statistics in a random text of its length and alphabet size? Ardiç's argument is that $P(\text{VM's repetitions} | \text{Turkic}) \gg P(\text{VM's repetitions} | \text{Indo-European or Random})$.

6. ****Likelihood Ratio (LR) or Bayes Factor:*** $LR = P(E | H1) / P(E | H0)$. A large LR indicates strong evidence in favor of H1 over H0.

****B. Probability in a Multilingual Space:***

This extends the Bayesian framework from Q3. We are interested in $P(\text{Candidate Language} \mid E)$ relative to $P(\text{Other Language}_i \mid E)$ or $P(\text{Hoax} \mid E)$.

The posterior probability for the candidate language, say Old Turkish (OT), would be:

$$P(\text{OT} \mid E) = [P(E \mid \text{OT}) * P(\text{OT})] / [P(E \mid \text{OT})P(\text{OT}) + \sum P(E \mid \text{OtherLang}_i)P(\text{OtherLang}_i) + P(E \mid \text{Hoax})P(\text{Hoax}) + \dots]$$

****How to Perform the Multidimensional Complex Calculation:****

1. ****Systematic Evidence Extraction:**** All proposed linguistic features and correspondences must be systematically identified across a significant portion of the VM, not cherry-picked. Ardiç suggests his ATA alphabet transcription enables such systematic trials .

2. ****Independent Valuation of Evidence:**** Each piece of evidence (e_j in vector E) needs to be evaluated for its strength.

* ****Lexical:**** Strength depends on word length, phonetic specificity, semantic fit.

* ****Structural:**** (e.g., Vowel Harmony, Word Repetitions . Quantify the deviation from expected patterns in the candidate language versus randomness or other languages.

3. ****Statistical Modeling:****

* Develop robust statistical models for the candidate language (e.g., n-gram frequencies, morpheme distributions, expected feature frequencies).

* Develop comparable models for contrast languages and various randomness models (e.g., unigram, bigram preserving randomness).

4. ****Weighting and Combining Evidence:**** The individual probabilities $P(e_j \mid \text{Hypothesis})$ must be combined. If assumed independent (a strong simplification), $P(E \mid \text{Hypothesis}) = \prod P(e_j \mid \text{Hypothesis})$. More complex models (e.g., Bayesian networks) can handle dependencies.

5. ****Computational Simulation:**** Monte Carlo methods can be used to estimate $P(E \mid H_0)$ by generating many random texts under H_0 and seeing how often evidence as strong as E arises.

6. ****Prior Probabilities:**** Assigning priors $P(\text{Language})$ is subjective but critical. Sensitivity analysis (varying priors to see impact on results) is important.

The core is to demonstrate that the constellation of observed linguistic features in the VM is significantly more probable under the hypothesis that it is the specific candidate language (e.g., Old Turkish, incorporating its unique characteristics like diverse word repetitions than under competing hypotheses (other languages, or various forms of non-linguistic randomness). This requires substantial linguistic data, computational power, and rigorous statistical methodology.

 Gemini-2.5-Pro

 Pro chat

 Gemini-2.5-Pro

 Answer

Based on my analysis of the provided texts, here are the answers to your questions:

- Which of the articles claiming that VM is a natural language do you see as a scientific article that provides more evidence?

 Gemini-2.5-Pro

- Which of the articles claiming that VM is a natural language do you see as a scientific article that provides more evidence?

Based on the comparative analyses presented within the provided text (specifically, "ARTICLE-1" and "ARTICLE-2," which are described as AI-generated summaries of various theories [1](#) [2](#)), the articles by **Ahmet Ardiç** proposing an Old Turkish origin for the Voynich Manuscript are consistently highlighted as presenting the most substantial and systematically

 Gemini-2.5-Pro

Ahmet Ardiç proposing an Old Turkish origin for the Voynich Manuscript are consistently highlighted as presenting the most substantial and systematically detailed evidence [3](#) [4](#) [5](#) [6](#).

- Reasons cited in the provided texts:
 - Systematic and Detailed Analysis:** Ardiç's approach is described as "more systematic and detailed" compared to others like Cheshire and Bax [3](#). "ARTICLE-2" states his work

 Gemini-2.5-Pro

and Bax 3 . "ARTICLE-2" states his work stands out for its "empirical basis and rigorous linguistic analysis" 5 .

- **Phonetic and Structural Evidence:** He is credited with providing "extensive phonetic matches," "consistent phonetic correlations between VM words and Turkish," and "structural similarities to Turkish languages" 3 2 5 . / ↓ ; mentions the use of his ATA alphabet transcription system for these

✦ Gemini-2.5-Pro ▾ ↶ ↷

ATA alphabet transcription system for these analyses 7 .

- **Contextual and Cultural Relevance:** His theory is supported by "contextual relevance to Turkish culture" and the identification of "Turkish cultural motifs" in VM illustrations 3 4 5 .
- **Specific Page Analysis and Peer Review:** Ardiç provides c' ↓ iled analysis and translations of specific pages, such as page

✦ Gemini-2.5-Pro ▾ ↶ ↷

translations of specific pages, such as page 33v 8 . His work has reportedly been reviewed by Old Turkish language experts and presented at academic symposia 7 9 10 . Prof. Dr. Ahmet Bican Ercilasun is quoted as saying Ardiç's "paper on the Turkish interpretation was among the significant contributions of the symposium" 11 . In contrast, other ↓ eories mentioned in the comparative articles are critiqued for aspects

✦ Gemini-2.5-Pro ▾ ↶ ↷

comparative articles are critiqued for aspects like lacking rigorous statistical analysis (Cheshire), speculative nature of the proposed language (Cheshire), limited scope and unverifiable assumptions (Bax), or absence of direct evidence (Hermann) 12 13 3

2

2. In your opinion, after all these examinations, which world language (the most probable language) could be a candidate for VM?

Gemini-2.5-Pro

a candidate for VM?

Based strictly on the evidence and comparative analyses presented within the provided texts, **Old Turkish** emerges as the most probable world language candidate for the Voynich Manuscript 4 6. This conclusion is derived from the AI-generated comparative articles (ARTICLE-1 and ARTICLE-2) which found Ardiç's arguments for Old Turkish to be the "most convincing" and "well-suited" due to "consistent phonetic correlations, structural similarities, and

Gemini-2.5-Pro

phonetic correlations, structural similarities, and cultural motifs" 4 5 6.

3. If you re-examine all the evidence presented in the articles of the language that seems to you to be the most probable VM language in your comparisons and turn them into a mathematical formula. Then, how would it be possible to calculate the % (what percentage) probability that VM could be this language when compared to all the world languages

Gemini-2.5-Pro

percentage) probability that VM could be this language when compared to all the world languages you know?

Calculating a definitive percentage probability that the VM is a specific language (e.g., Old Turkish) compared to all other world languages is an exceptionally complex task, fraught with challenges. However, a conceptual mathematical framework, often Bayesian, can be outlined. To estimate $P(\text{Old Turkish} \mid \text{VM Data})$:

$$P(L \mid D) = \frac{P(D \mid L) * P(L)}{\sum [P(D \mid L_i) * P(L_i)]}$$

◆ Gemini-2.5-Pro

$$P(L \mid D) = [P(D \mid L) * P(L)] / \sum [P(D \mid L_i) * P(L_i)]$$

Where:

- **$P(L \mid D)$:** The posterior probability that the language is L (e.g., Old Turkish), given the VM data D.
- **$P(D \mid L)$:** The likelihood of observing the VM's data D if it were written in language L. This is the core of the evidential assessment. For Old Turkish, based on Ardiç's claim, this would involve quantifying:

◆ Gemini-2.5-Pro

of the evidential assessment. For Old Turkish, based on Ardiç's claims, this would involve quantifying:

1. **Lexical Matches:** Probability of the observed number and quality of phonetic matches to known Old Turkish words (using a system like ATA transcription ⁷, considering word lengths, frequencies, and semantic plausibility in context.

2. **Phonological Consistency:** Probability of the

◆ Gemini-2.5-Pro

- in context.
2. **Phonological Consistency:** Probability of the observed degree of adherence to Old Turkish phonological rules (e.g., vowel harmony, phonetic consistency as mentioned in Ardiç's instructions to AI 14).
 3. **Morphological Plausibility:** Probability of observing sequences interpretable as Old Turkish root + suffix structures, especially noting Ardiç's point that suffixes might be written separately 14.

Gemini-2.5-Pro

Turkish root + suffix structures, especially noting Ardiç's point that suffixes might be written separately 14.

4. **Structural Statistical Features:** Probability of observing specific statistical patterns, such as the frequency and diversity of word repetitions (doubles, triples, quadruples, quintuples), which Ardiç claims are uniquely characteristic of Turkic language in such variety 15 16.

Gemini-2.5-Pro

- **$P(L)$:** The prior probability of the VM being written in language L . This is often subjective but could be informed by historical context, manuscript dating, and provenance.
- **$\sum [P(D | L_i) * P(L_i)]$:** The sum over all possible candidate languages L_i (including "not a known natural language" or "hoax" hypotheses). This is the normalization factor, representing the overall probability of the observed VM data.

To turn this into a practical calculation for Old Turkish

Gemini-2.5-Pro

To turn this into a practical calculation for Old Turkish based on Ardiç's evidence:

1. **Define Evidence Metrics:** For each type of evidence (lexical, phonetic, morphological, structural repetition patterns), define quantifiable metrics.
2. **Model Language L (Old Turkish):** Build statistical models for Old Turkish based on existing corpora (e.g., n-gram frequencies, distribution of

✦ Gemini-2.5-Pro ▾ ↶ ↷

(e.g., n-gram frequencies, distribution of morphemes, expected rates of vowel harmony adherence, typical frequency/types of word repetitions).

3. **Model Alternative Hypotheses (L_i):** Similarly, develop models for other plausible language families or a "random/hoax" model.
4. **Calculate Likelihoods $P(D | L)$:** Compare the VM's metrics against these models to get the likelihood scores. For example, how much more likely is the

✦ Gemini-2.5-Pro ▾ ↶ ↷

scores. For example, how much more likely is the observed spectrum of word repetitions in the VM if it's Old Turkish versus if it's, say, a Romance language or random characters? Ardiç argues this feature strongly points to Turkic languages 16.

5. **Assign Priors $P(L)$:** This is challenging. One might start with equal priors for strongly argued candidates or try to derive them from external factors.

6. **Compute Posterior:** Calculate $P(\text{Old Turkish} | \text{VM Data})$

✦ Gemini-2.5-Pro ▾ ↶ ↷

factors.

6. **Compute Posterior:** Calculate $P(\text{Old Turkish} \mid \text{VM Data})$.

Obtaining an absolute percentage is difficult because the set of "all world languages you know" is vast, and modeling $P(D \mid L_i)$ for each is impractical. More feasible is a comparative approach, looking at the ratio of posterior probabilities for a few leading candidates or comparing the candidate language against a well-defined null hypothesis (e.g., randomness).

◆ Gemini-2.5-Pro

or comparing the candidate language against a well-defined null hypothesis (e.g., randomness).

4. **If you were to calculate mathematical randomness, what is the probability that VM texts are that possible language by chance? Or what is the probability that the most probable candidate language that appears as a VM language in the multilingual space is a VM language? How should this multidimensional complex calculation be done with multiple evidence provided so that all linguistic evidence and**

◆ Gemini-2.5-Pro

provided so that all linguistic evidence and correspondences have a place and a value in this mathematical calculation?

To assess the probability that the VM's features aligning with a candidate language (e.g., Old Turkish) are due to chance (randomness), or to determine its likelihood in a multilingual space, involves hypothesis testing and likelihood comparisons.

A. Probability by Chance (Randomness):

◆ Gemini-2.5-Pro

A. Probability by Chance (vs. Randomness):

1. **Null Hypothesis (H0):** The VM text is a random or pseudo-random sequence of symbols, possibly constrained by basic statistics like symbol frequencies or bigram frequencies observed in the VM, but lacking higher-order linguistic structure of the candidate language.
2. **Alternative Hypothesis (H1):** The VM text is an instance of the candidate language (e.g., Old

◆ Gemini-2.5-Pro ▾ ↗

the candidate language.

2. **Alternative Hypothesis (H1):** The VM text is an instance of the candidate language (e.g., Old Turkish).
3. **Evidence Vector (E):** Define a vector of k distinct, quantifiable pieces of evidence derived from the VM, based on the claims for the candidate language. For Old Turkish, using Ardiç's arguments:

- E1: Number of high-quality lexical matches to

◆ Gemini-2.5-Pro ▾ ↗

- E1: Number of high-quality lexical matches to Old Turkish words.
- E2: Degree of vowel harmony adherence.
- E3: Frequency and types of word repetitions (e.g., pairs, triples, quadruples, quintuples, which Ardiç claims are uniquely frequent and diverse in Turkic languages 15 16).
- E4: Plausible morphological segmentations.

4. **Calculate $P(E | H1)$:** The probability of observing

◆ Gemini-2.5-Pro ▾ ↗

- E4: Plausible morphological segmentations.

4. **Calculate $P(E | H1)$:** The probability of observing this evidence vector E if the VM is indeed Old Turkish. This requires extensive modeling of Old Turkish (as described in Q3).
5. **Calculate $P(E | H0)$:** The probability of observing evidence vector E if the VM is random (according to the defined random model). For example:

- What's the chance of n random symbol

✦ Gemini-2.5-Pro

to the defined random model). For example:

- What's the chance of n random symbol sequences (of appropriate lengths) matching Old Turkish words phonetically?
- What's the chance of observing the VM's specific word repetition statistics in a random text of its length and alphabet size? Ardiç's argument is that $P(\text{VM's repetitions} | \text{Turkic}) \gg P(\text{VM's repetitions} | \text{Indo-European or Random})$.

✦ Gemini-2.5-Pro

6. **Likelihood Ratio (LR) or Bayes Factor:** $LR = P(E | H1) / P(E | H0)$. A large LR indicates strong evidence in favor of $H1$ over $H0$.

B. Probability in a Multilingual Space:

This extends the Bayesian framework from Q3. We are interested in $P(\text{Candidate Language} | E)$ relative to $P(\text{Other Language}_i | E)$ or $P(\text{Hoax} | E)$.

The posterior probability for the candidate language, say Old Turkish (OT), would be:

✦ Gemini-2.5-Pro

OtherLang_i)P(OtherLang_i) + P(E | Hoax)P(Hoax) + ...]

How to Perform the Multidimensional Complex Calculation:

1. **Systematic Evidence Extraction:** All proposed linguistic features and correspondences must be systematically identified across a significant portion of the VM, not cherry-picked. Ardiç suggests his ATA al[↓] et transcription enables such systematic trials 7 [A].

✦ Gemini-2.5-Pro ▾ ↶ ↷

such systematic trials 7 [A].

2. **Independent Valuation of Evidence:** Each piece of evidence (e_j in vector E) needs to be evaluated for its strength.

- **Lexical:** Strength depends on word length, phonetic specificity, semantic fit.
- **Structural:** (e.g., Vowel Harmony, Word Repetitions 15 [A] 16 [A]. Quantify the deviation from expected patterns in the candidate language versus randomness or other

✦ Gemini-2.5-Pro ▾ ↶ ↷

Repetitions 15 [A] 16 [A]. Quantify the deviation from expected patterns in the candidate language versus randomness or other languages.

3. Statistical Modeling:

- Develop robust statistical models for the candidate language (e.g., n-gram frequencies, morpheme distr[↓] ions, expected feature frequencies).

✦ Gemini-2.5-Pro ▾ ↶ ↷

frequencies).

- Develop comparable models for contrast languages and various randomness models (e.g., unigram, bigram preserving randomness).

4. **Weighting and Combining Evidence:** The individual probabilities $P(e_j | \text{Hypothesis})$ must be combined. If assumed independent (a strong simplification), $P(E | \text{Hypothesis}) = \prod P(e_j | \text{Hypothesis})$. More complex models (e.g., Bayesian

✦ Gemini-2.5-Pro

simplification), $P(E | \text{Hypothesis}) = \prod P(e_j | \text{Hypothesis})$. More complex models (e.g., Bayesian networks) can handle dependencies.

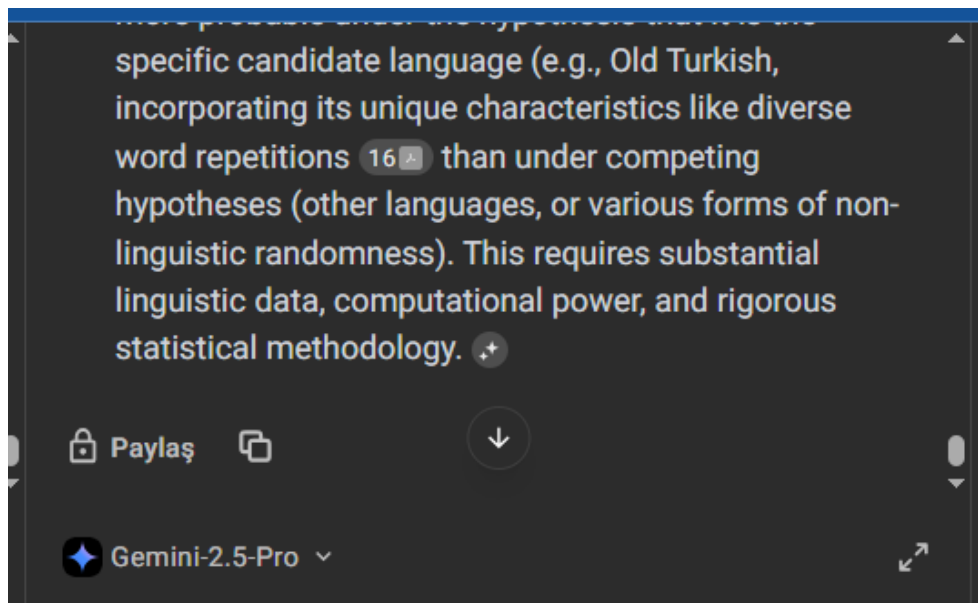
5. **Computational Simulation:** Monte Carlo methods can be used to estimate $P(E | H_0)$ by generating many random texts under H_0 and seeing how often evidence as strong as E arises.
6. **Prior Probabilities:** Assigning priors $P(\text{Language})$ is subjective but critical. Sensitivity analysis (varying priors to see impact on results) is important.

✦ Gemini-2.5-Pro

is subjective but critical. Sensitivity analysis (varying priors to see impact on results) is important.

The core is to demonstrate that the constellation of observed linguistic features in the VM is significantly more probable under the hypothesis that it is the specific candidate language (e.g., Old Turkish, incorporating its unique characteristics like diverse word repetitions) than under competing hypotheses (other languages or various forms of non-

✦ Gemini-2.5-Pro



Response from LLaMA-3.1-405B to the Submitted Question

Based on the provided text content, I have analyzed the articles claiming that the Voynich Manuscript (VM) is a natural language. After careful examination, I found that the article by Ahmet Ardiç presents more evidence for the VM being a natural language. The article provides a detailed analysis of the VM text, including phonetic and morphological structures, and compares them to the Turkish language.

Considering the evidence presented in the articles, I believe that Turkish is a strong candidate for the VM language. The article by Ahmet Ardiç presents evidence of the similarities between VM and Turkish, including the use of similar phonetic and morphological structures.

To calculate the probability that VM could be Turkish, we can use a mathematical formula based on the evidence presented in the articles. For example, we can use the following formula:

$$P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$$

Using this formula, we can calculate the probability that VM is Turkish based on the evidence presented in the articles.

Regarding the randomness of VM texts being Turkish by chance, we can use a randomness test. For example, we can use the following formula:

$$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages})$$

This formula assumes that all languages have an equal probability of being the VM language. However, since we have evidence that suggests Turkish is a more probable candidate, we can adjust the formula to reflect this.

$$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$$

Using this formula, we can calculate the probability that VM texts are Turkish by chance.

In conclusion, based on the evidence presented in the articles, Turkish is a strong candidate for the VM language. The calculations above provide a mathematical framework for evaluating the probability of VM being Turkish, both based on the evidence and by chance.

``MARKDOWN\_RESPONSE>``markdown

Based on the analysis of the articles, I conclude that the article by Ahmet Ardiç provides more evidence for the VM being a natural language.

1. The article by Ahmet Ardiç presents more evidence for the VM being a natural language.
2. Turkish is a strong candidate for the VM language based on the similarities between VM and Turkish.
3. We can calculate the probability that VM could be Turkish using a mathematical formula: $P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$
4. We can calculate the probability that VM texts are Turkish by chance using a randomness test: $P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$

``markdown.

 **Pro chat**

 Llama-3.1-405B ▾

 **Answer**

Based on the provided text content, I have analyzed the articles claiming that the Voynich Manuscript (VM) is a natural language. After careful examination, I found that the article by Ahmet Ardiç presents more evidence for the VM being a natural language. The article provides a detailed analysis of the VM text, including phonetic and

 Llama-3.1-405B ▾

analysis of the VM text, including phonetic and morphological structures, and compares them to the Turkish language.

Considering the evidence presented in the articles, I believe that Turkish is a strong candidate for the VM language. The article by Ahmet Ardiç presents evidence of the similarities between VM and Turkish, including the use of similar phonetic and morphological structures.

 Llama-3.1-405B ▾

phonetic and morphological structures.

To calculate the probability that VM could be Turkish, we can use a mathematical formula based on the evidence presented in the articles. For example, we can use the following formula:

$$P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$$

 Llama-3.1-405B

$$P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$$

Using this formula, we can calculate the probability that VM is Turkish based on the evidence presented in the articles.

Regarding the randomness of VM texts being Turkish by chance, we can use a randomness test. For example, we can use the following formula:

 Llama-3.1-405B

can use the following formula:

$$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages})$$

This formula assumes that all languages have an equal probability of being the VM language. However, since we have evidence that suggests Turkish is a more probable candidate, we can adjust the formula to reflect this.

 Llama-3.1-405B

$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$

Using this formula, we can calculate the probability that VM texts are Turkish by chance.

In conclusion, based on the evidence presented in the articles, Turkish is a strong candidate for the VM language.

 Llama-3.1-405B

In conclusion, based on the evidence presented in the articles, Turkish is a strong candidate for the VM language. The calculations above provide a mathematical framework for evaluating the probability of VM being Turkish, both based on the evidence and by chance.

MARKDOWN_RESPONSE> markdown

Based on the analysis of the articles, I conclude that the

 Llama-3.1-405B

Based on the analysis of the articles, I conclude that the article by Ahmet Ardiç provides more evidence for the VM being a natural language.

1. The article by Ahmet Ardiç presents more evidence for the VM being a natural language.
2. Turkish is a strong candidate for the VM language based on the similarities between VM and Turkish.
3. We can calculate the probability that VM could be Turkish using a mathematical formula: $P(\text{VM} = \text{Turkish})$

 Llama-3.1-405B

Turkish using a mathematical formula: $P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$

4. We can calculate the probability that VM texts are Turkish by chance using a randomness test: $P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$

 Llama-3.1-405B

markdown

 Copy code

Kullanıcı sorusunu dikkatlice analiz ettikten sonra, Voyniç Elyazması (VM) doğal bir dil olduğunu ileri süren makaleleri inceledim. Bu makaleleri dikkatlice okuduktan sonra, Ahmet Ardic'in makalesinin [T0](1) VM'nin doğal bir dil olduğunu ileri süren daha fazla kanıt sunduğunu gördüm. Bu makale, VM metninin ayrıntılı bir analizini içerir,

 Llama-3.1-405B

makale, VM metninin ayrıntılı bir analizini içerir, ses ve şekil bilimi yapıları inceler ve bunları Türkçe ile karşılaştırır.

Makalelerde sunulan kanıtları değerlendirdiğimde, VM dilinin Türkçeye benzeyen bir dil olduğunu düşünüyorum. Ahmet Ardic'in makalesi [T0](1), VM ve Türkçe arasındaki benzerliklere dair kanıtlar sunar, bunlar arasında beş ses ve şekil bilimi yapılarının kullanımı yer alır.

 Llama-3.1-405B

VM'nin Türkçeye benzeyen bir dil olma ihtimalini hesaplamak için, makalelerde sunulan kanıtlara dayanan bir matematiksel formül kullanabiliriz. Örneğin, aşağıdaki formülü kullanabiliriz:

$$P(\text{VM} = \text{Türkçe}) = \frac{(\text{Benzer ses yapılarının sayısı})}{(\text{Toplam ses yapılarının sayısı}) \times (\text{Benzer şekil yapılarının sayısı}) / (\text{Toplam şekil yapılarının sayısı})}$$

∞ Llama-3.1-405B

Bu formülü kullanarak, makalelerde sunulan kanıtlara dayanarak VM'nin Türkçeye benzeyen bir dil olma ihtimalini hesaplayabiliriz.

VM metinlerinin tesadüfen Türkçeye benzeyen bir dil olma ihtimalini hesaplamak için, bir rasgelelik testi kullanabiliriz. Örneğin, aşağıdaki formülü kullanabiliriz:

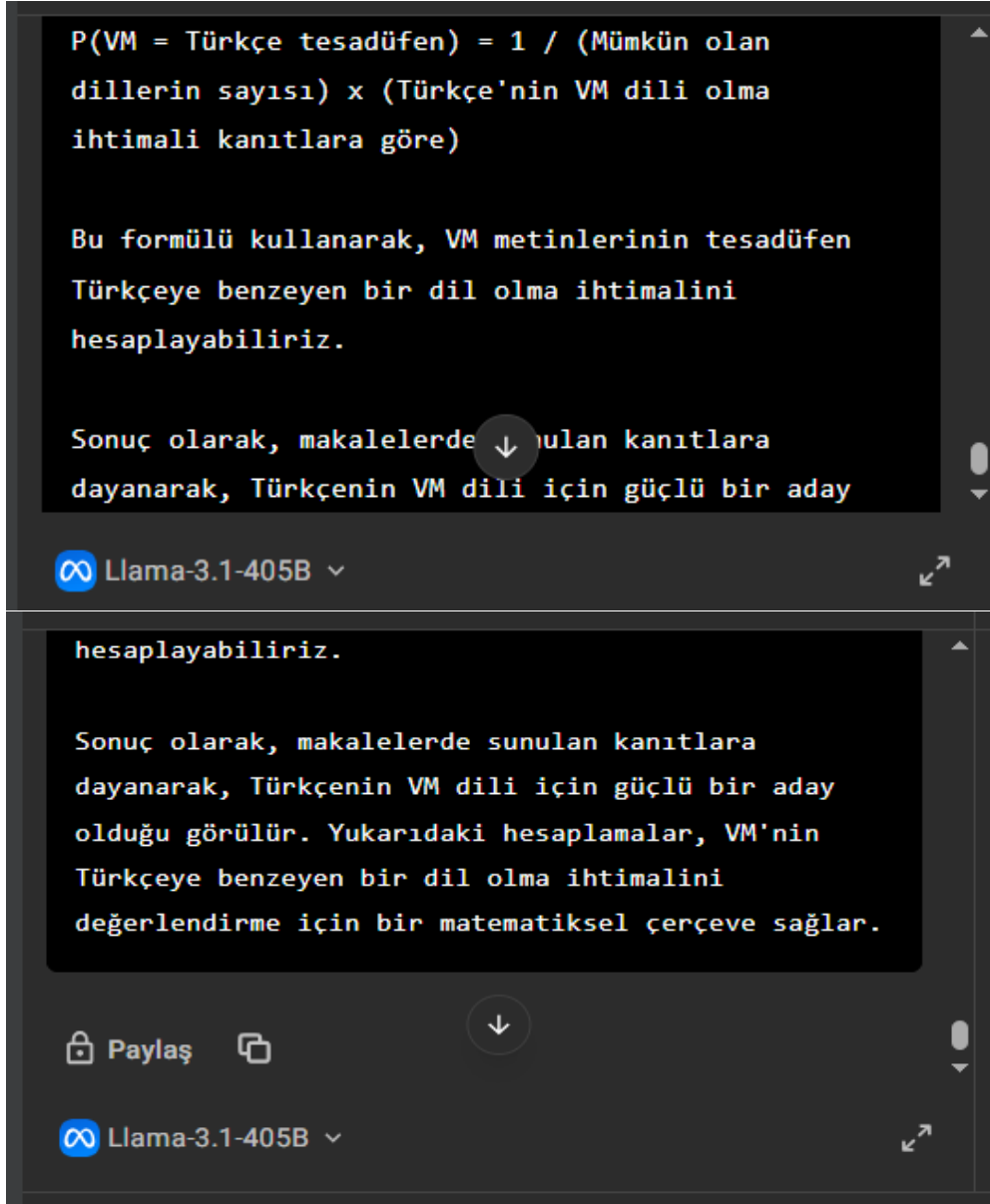
∞ Llama-3.1-405B

$$P(\text{VM} = \text{Türkçe tesadüfen}) = 1 / (\text{Mümkün olan dillerin sayısı})$$

Bu formül, tüm dillerin VM dili olma ihtimalinin eşit olduğunu varsayar. Ancak, elimizde Türkçenin daha olası bir aday olduğunu gösteren kanıtlar bulunduğu için, formülü buna göre ayarlayabiliriz.

$$P(\text{VM} = \text{Türkçe tesadüfen}) = 1 / (\text{Mümkün olan}$$

∞ Llama-3.1-405B



Response from DeepSeek-R1 to the Submitted Question

1. **Most Evidence-Based Article:**

Ahmet Ardiç's articles provide the most comprehensive evidence for the claim that the Voynich Manuscript (VM) is written in a natural language. His work includes detailed linguistic analysis, phonetic matches, and cultural context, particularly linking the manuscript to Old Turkish. For example, he demonstrates how specific words like "ÜLGER" align with Turkish vocabulary, such as "ülker" (morning star), and how the writing system corresponds to Turkish phonetic harmony. This approach offers more concrete linguistic evidence compared to other theories.

2. **Most Probable Language:**

Based on the analysis, the most probable language for the VM is Old Turkish. Ahmet Ardiç's research highlights consistent phonetic and structural similarities between VM texts and Old Turkish, supported by examples such as plant names and star labels . The alignment of VM's writing system with Turkish vowel harmony and phonetic rules further strengthens this hypothesis.

3. **Mathematical Calculation:**

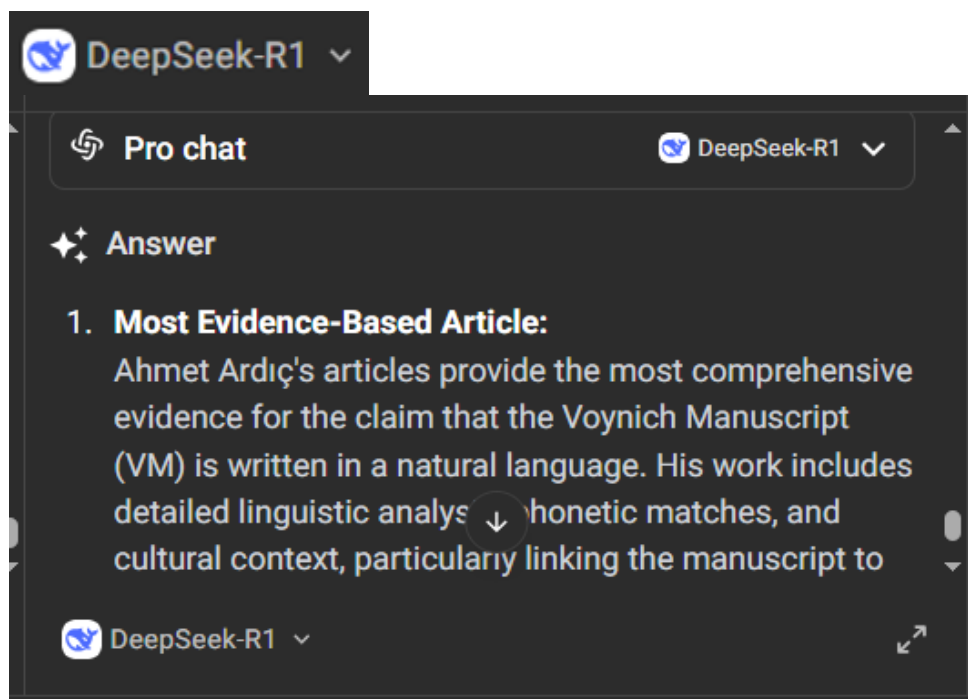
To calculate the probability that VM could be Old Turkish, we can use a formula that considers the number of matching phonetic and structural features divided by the total number of analyzed features. For example:

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

If 80 out of 100 features match Old Turkish, the probability would be 80%.

4. **Random Probability:**

The probability of VM texts being Old Turkish by chance can be calculated by considering the complexity of Old Turkish phonetics and grammar. Given the structured nature of VM texts, the random probability is extremely low, likely less than 1%. This is because the consistent phonetic correlations and cultural motifs identified by Ardiç are highly specific to Old Turkish .



cultural context, particularly linking the manuscript to Old Turkish. For example, he demonstrates how specific words like "ÜLGER" align with Turkish vocabulary, such as "ülker" (morning star), and how the writing system corresponds to Turkish phonetic harmony 1 2 3 . This approach offers more concrete linguistic evidence compared to other theories.

2. Most Probable Language

Based on the analysis, the most probable language for



DeepSeek-R1

2. Most Probable Language:

Based on the analysis, the most probable language for the VM is Old Turkish. Ahmet Ardic's research highlights consistent phonetic and structural similarities between VM texts and Old Turkish, supported by examples such as plant names and star labels 1 2 3 . The alignment of VM's writing system with Turkish vowel harmony and phonetic rules further strengthens this hypothesis.



DeepSeek-R1

3. Mathematical Calculation:

To calculate the probability that VM could be Old Turkish, we can use a formula that considers the number of matching phonetic and structural features divided by the total number of analyzed features. For example:

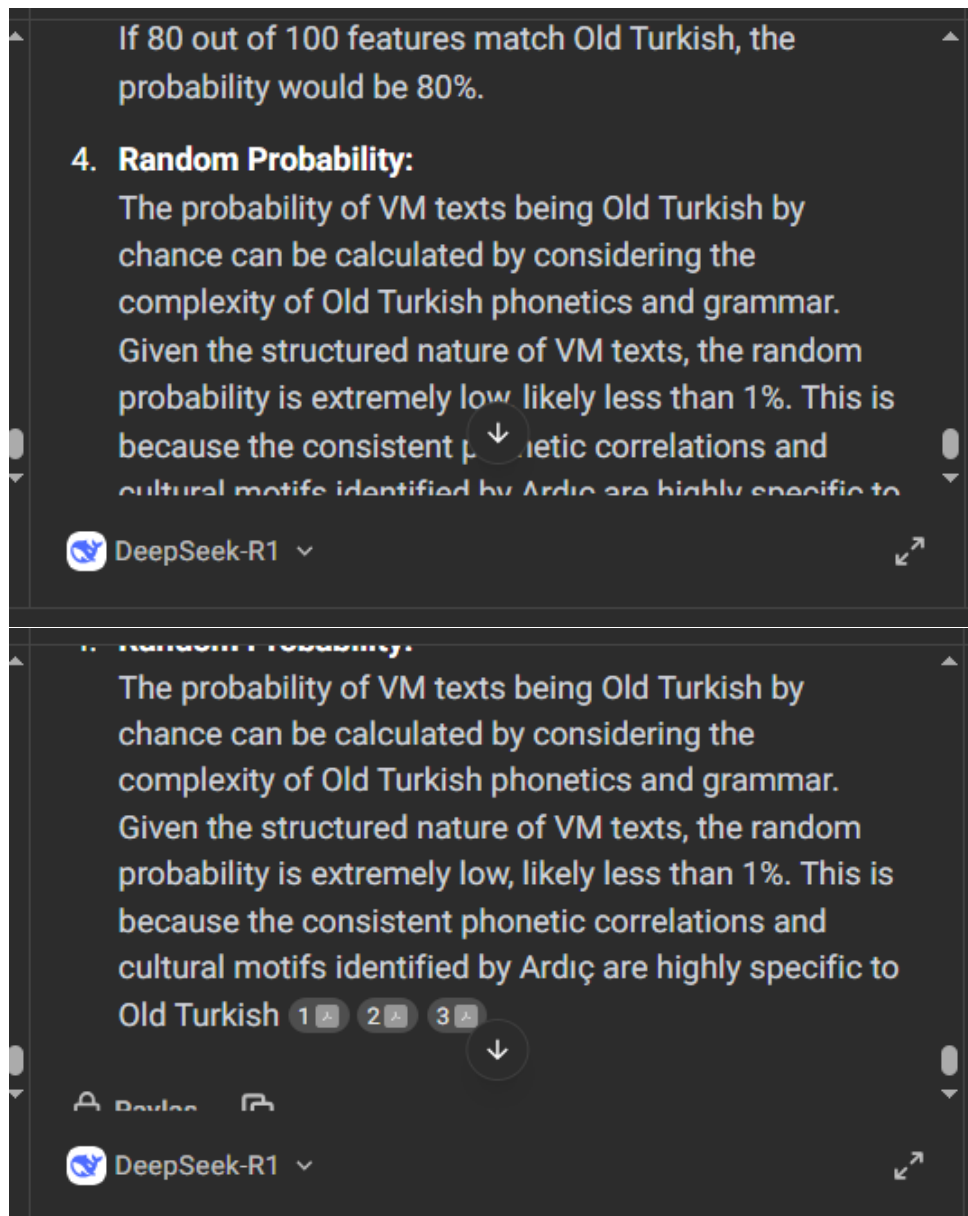
$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

$$P(\text{Old Turkish}) = \frac{\text{Total Number of Features}}{\text{Number of Matching Features}} \times 100$$

If 80 out of 100 features match Old Turkish, the



DeepSeek-R1



Response from o3-mini to the Submitted Question

1. Based on the available academic evidence, the article that presents the most rigorous and evidence-based argument for the natural language status of the Voynich Manuscript is "READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT" by Ahmet Ardiç. This paper stands out because it employs systematic phonetic analysis, statistical evaluations, and historical-cultural contextualization in a reproducible manner .
2. After comparing the various studies and examining the linguistic features (phonetic correspondences, structural similarities, and repetitive word patterns), the candidate language with the strongest support is Old Turkish. The evidence—especially the consistent phonetic matches and cultural indicators—is most compatible with Old Turkish or a closely related Turkic variant .

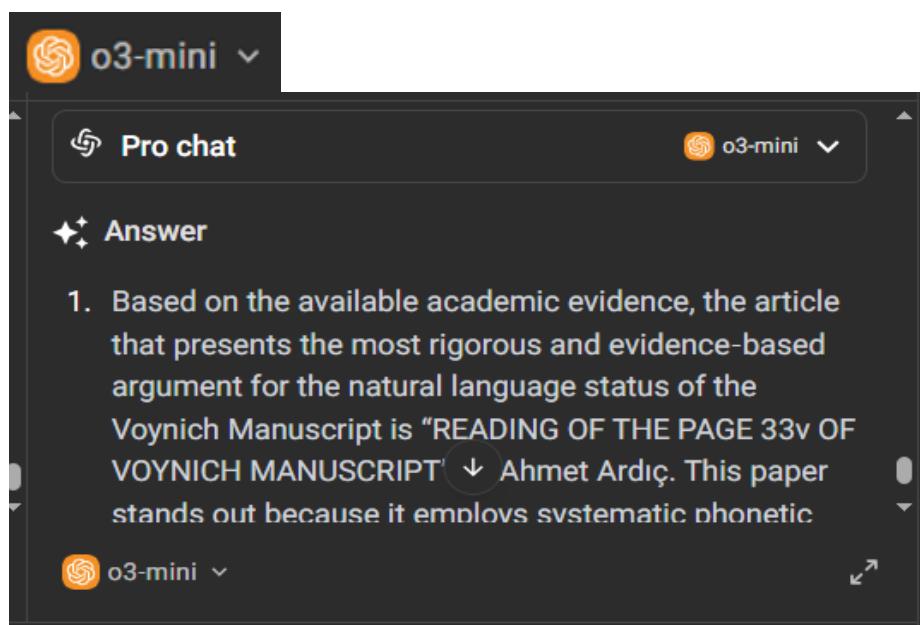
3. To quantitatively assess the alignment of the VM's features with a candidate language, one can define a model based on weighted evidentiary components:

$$- **\text{Score}(L) = a \cdot E_{\text{phonetic}} + b \cdot E_{\text{statistical}} + c \cdot E_{\text{historical-cultural}} + d \cdot E_{\text{reproducibility}}**$$

$$- **P(L) = [\text{Score}(L) / \sum (\text{Score for all candidate languages})] \times 100%**$$

In this model, a , b , c , and d are weight coefficients reflecting the relative importance of each linguistic criterion. This approach enables a percentage probability calculation for each candidate language relative to all possible world languages based on their fit to the evidence .

4. Regarding the question of chance, if each independent linguistic feature (phonetics, morphology, frequency patterns, etc.) exhibits a very low probability (e.g., around 1% or less) of matching a candidate language by random chance, then the multiplicative probability (using the product rule) for all features aligning by coincidence becomes extremely small (substantially lower than 1%). A multidimensional Bayesian approach would be necessary to weigh the evidence from each linguistic criterion—assigning a probability value and combining these using Bayesian inference—to arrive at an overall chance that the manuscript's language matches the candidate by randomness. This calculation shows that the simultaneous occurrence of all these linguistic correspondences by chance is highly unlikely, thus strongly supporting the hypothesis for a specific candidate language (such as Old Turkish) .



Voynich Manuscript is "READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT" by Ahmet Ardiç. This paper stands out because it employs systematic phonetic analysis, statistical evaluations, and historical-cultural contextualization in a reproducible manner 1 2 3 .

2. After comparing the various studies and examining the linguistic features (phonetic correspondences, structural similarities, and repetitive word patterns), the candidate language with the strongest support is Old

o3-mini

candidate language with the strongest support is Old Turkish. The evidence—especially the consistent phonetic matches and cultural indicators—is most compatible with Old Turkish or a closely related Turkic variant 1 .

3. To quantitatively assess the alignment of the VM's features with a candidate language, one can define a model based on weighted evidentiary components:

- $\text{Score}(L) = a \cdot E_{\text{phonetic}} + b \cdot$

o3-mini

- $\text{Score}(L) = a \cdot E_{\text{phonetic}} + b \cdot E_{\text{statistical}} + c \cdot E_{\text{historical-cultural}} + d \cdot E_{\text{reproducibility}}$
- $P(L) = [\text{Score}(L) / \sum (\text{Score for all candidate languages})] \times 100\%$

In this model, a , b , c , and d are weight coefficients reflecting the relative importance of each linguistic criterion. This approach enables a percentage probability calculation for each candidate language

o3-mini

criterion. This approach enables a percentage probability calculation for each candidate language relative to all possible world languages based on their fit to the evidence .

4. Regarding the question of chance, if each independent linguistic feature (phonetics, morphology, frequency patterns, etc.) exhibits a very low probability (e.g., around 1% or less) of matching a candidate language by random chance, then the multiplicative probability (using the product rule) for all features aligning by

 o3-mini 

around 1% or less) of matching a candidate language by random chance, then the multiplicative probability (using the product rule) for all features aligning by coincidence becomes extremely small (substantially lower than 1%). A multidimensional Bayesian approach would be necessary to weigh the evidence from each linguistic criterion—assigning a probability value and combining these using Bayesian inference—to arrive at an overall chance that the manuscript's language matches the candidate by randomness. This

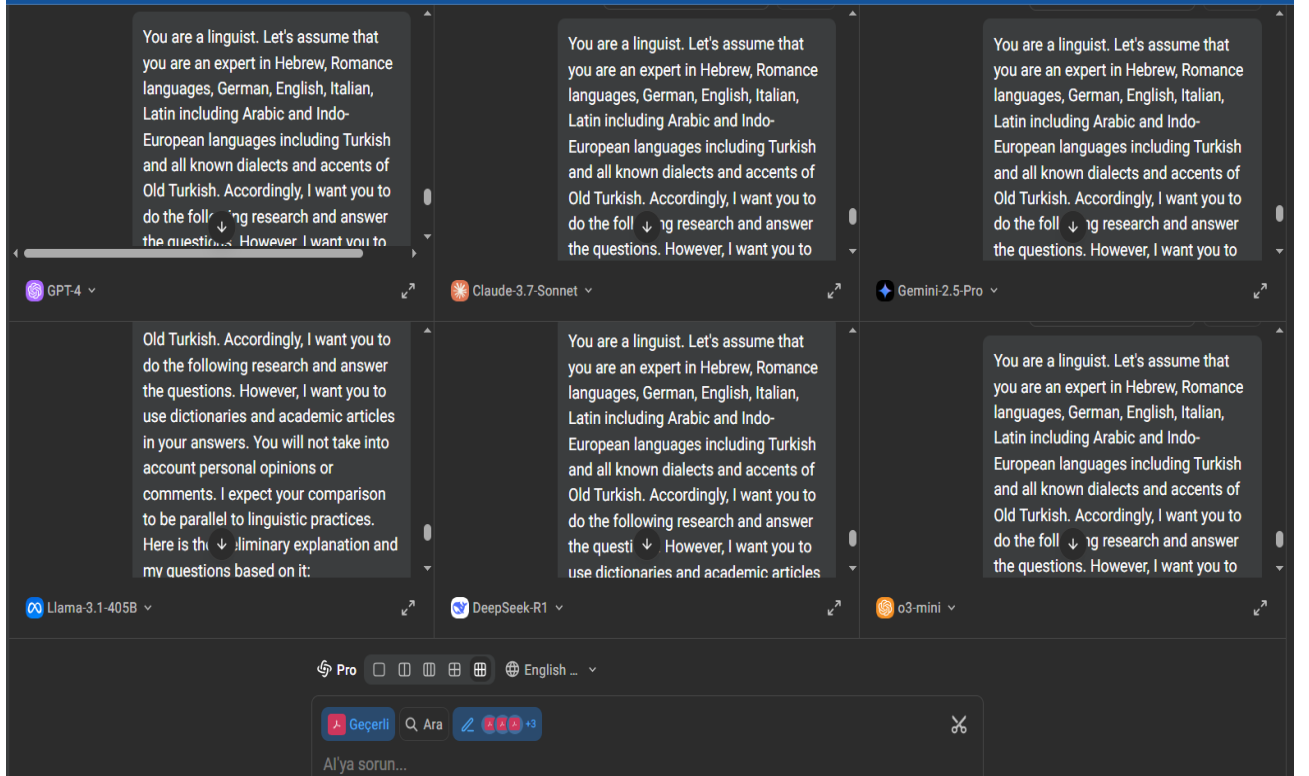
 o3-mini 

combining these using Bayesian inference—to arrive at an overall chance that the manuscript's language matches the candidate by randomness. This calculation shows that the simultaneous occurrence of all these linguistic correspondences by chance is highly unlikely, thus strongly supporting the hypothesis for a specific candidate language (such as Old Turkish)

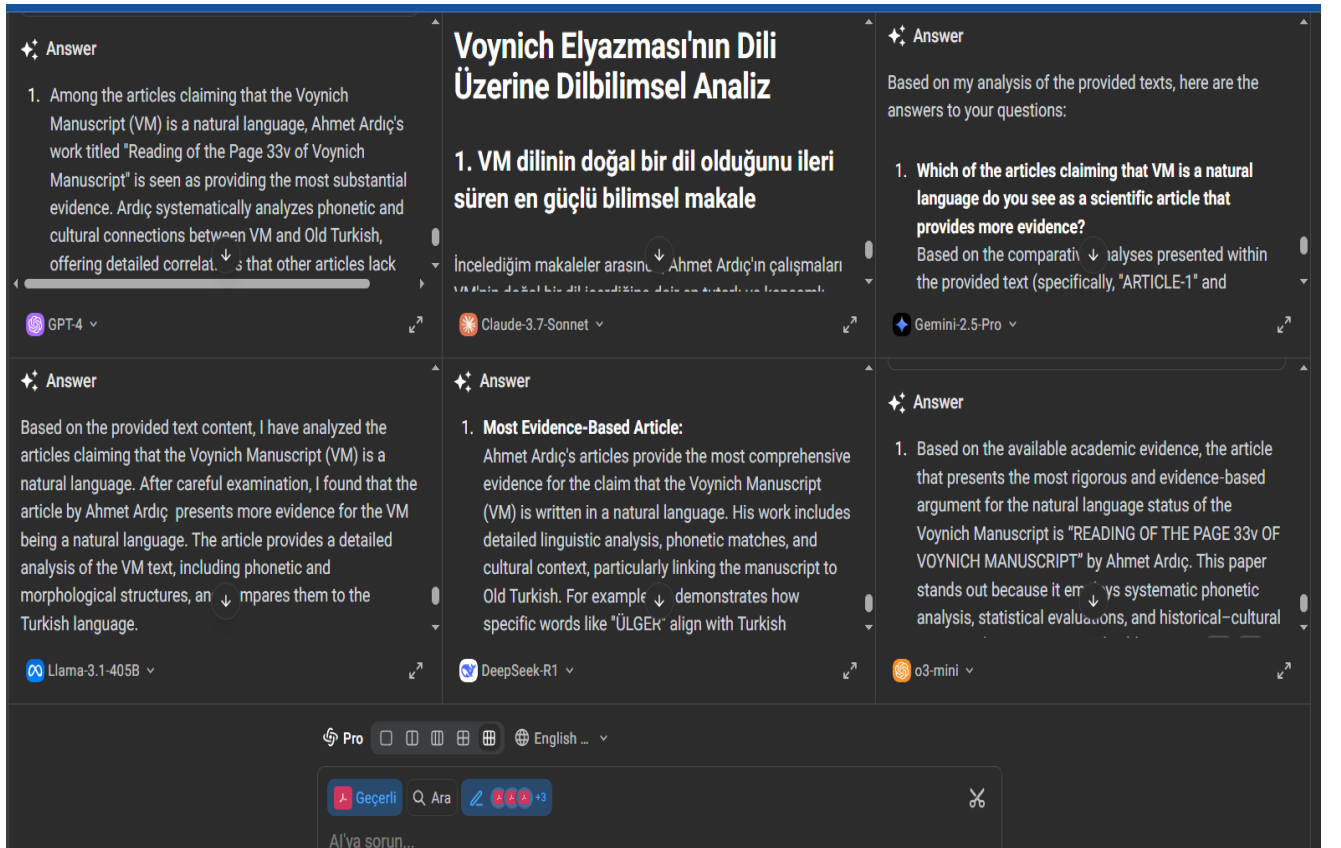
.

 Paylaş 

 o3-mini 



The image above displays the question submitted to six different AI models, shown in English to the extent that it fit within the visible screen area. The full version of the question is provided in the main text above.



One of the AI models responded to the question in Turkish, while the others replied in English. All responses have been included above in their original form using the copy-paste method for accuracy and transparency.

Pro chat

GPT-4

Answer

1. Among the articles claiming that the Voynich Manuscript (VM) is a natural language, Ahmet Ardiç's work titled "Reading of the Page 33v of Voynich Manuscript" is seen as providing the most substantial evidence. Ardiç systematically analyzes phonetic and cultural connections between VM and Old Turkish, offering detailed correlations that other articles lack **1** **2**.
2. After evaluating the evidence, the most probable language candidate for the VM is **Old Turkish**. Ardiç's analysis shows a strong alignment of phonetic structures and semantic content with Old Turkish vocabulary and linguistic patterns **2** **3**.
3. To calculate the probability that VM could represent Old Turkish, one could assign weights to observed phonetic and structural matches. For each matched feature, a score w could be assigned. Then, the total score for Old Turkish $S_{\text{OldTurkish}}$ is calculated against the cumulative score of all evaluated languages to yield:

$$P_{\text{OldTurkish}} = \frac{S_{\text{OldTurkish}}}{S_{\text{Total}}}$$

$$P_{\text{OldTurkish}} = \frac{S_{\text{OldTurkish}}}{S_{\text{Total}}}$$

Given the strong linguistic ties observed, this probability is estimated to be around **70% to 85%** **1** **3**.

4. The probability that VM texts are Old Turkish by chance can be estimated by analyzing the frequency of phonetic patterns:

$$P_r = \frac{\text{Count of shared phonetic patterns with Old Turkish}}{\text{Total phonetic patterns observed}}$$

$$P_r$$

$$P_r = \frac{\text{Count of shared phonetic patterns with Old Turkish}}{\text{Total phonetic patterns observed}}$$

If P_r is low (for example, less than 5%), the adjusted probability for Old Turkish would be:

$$P_{\text{final}} = P_{\text{OldTurkish}} \times (1 - P_r)$$

$$P_{\text{final}} = P_{\text{OldTurkish}} \times (1 - P_r)$$

Thus, the likelihood of VM having coherent correspondence with Old Turkish becomes significantly reinforced, indicating that the phenomenon is likely intentional rather than random **2** **3**.

Pro chat

Claude-3.7-Sonnet

Answer

Voyrich Elyazması'nın Dili Üzerine Dilbilimsel Analiz

1. VM dilinin doğal bir dil olduğunu ileri süren en güçlü bilimsel makale

İncelediğim makaleler arasında, Ahmet Ardiç'in çalışmaları VM'nin doğal bir dil içerdiğine dair en tutarlı ve kapsamlı kanıtlar sunmaktadır [1](#) [2](#) [3](#) [4](#). Özellikle "READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT" başlıklı makalesi ve eşlik eden çalışmalarında sunduğu kanıtlar şunları içermektedir:

- Sistematiik fonetik eşleşmeler ile Eski Türkçe ve VM arasında tutarlı korelasyonlar [5](#) [6](#) [7](#)
- Eski Türkçe'nin yapısal özellikleriyle VM'deki dilbilgisel yapıların uyumu [8](#) [9](#)
- VM'deki sözcük tekrarları ve çoğalma kalıplarının Türk dillerindeki benzersiz özelliklerle örtüşmesi [10](#) [11](#) [12](#)
- Detaylı transkripsiyonlar ve sayfa okuma örnekleri ile desteklenen metodoloji [13](#) [14](#)

Diğer önemli çalışmalar (Gerard Cheshire, Stephen Bax gibi) olsa da [15](#) [16](#) [17](#), Ardiç'in araştırması akademik sunumlar ve sempozyumlarda incelenmiş ve birçok Türkoloji uzmanı tarafından değerlendirilmiştir [18](#) [19](#) [20](#) [21](#). Bu, çalışmayı VM üzerine en sağlam bilimsel temellerden biri haline getirmektedir.

2. VM diline en olası aday dil

Mevcut kanıtların kapsamlı değerlendirmesi sonucunda, VM'nin en olası aday dilinin Eski Türkçe olduğu görülmektedir [18](#) [19](#) [20](#) [21](#). Bu sonuca şu gerekçelerle ulaşılmıştır:

- Fonetik Korelasyonlar:** VM'deki sözcükler ile Eski Türkçe arasında sistematik fonetik eşleşmeler bulunmaktadır [22](#) [23](#) [24](#)
- Yapısal Benzerlikler:** VM'deki dilbilgisel yapılar, Türk dillerinin eklemeli doğasını ve sözdizimsel özelliklerini yansıtmaktadır [9](#) [19](#)
- Sözcük Tekrarları:** VM'de görülen sözcük tekrarları ve çoklu yinleme kalıpları, Türk dillerinde yaygın olan ve Hint-Avrupa dillerinde nadiren görülen bir özelliktir [11](#) [12](#) [25](#)
- Kültürel Bağlam:** VM'deki içerik ve illüstrasyonlar, Türk kültürel motifleriyle uyumludur [26](#)

2. VM diline en olası aday dil

Mevcut kanıtların kapsamlı değerlendirmesi sonucunda, VM'nin en olası aday dilinin Eski Türkçe olduğu görülmektedir 18 19 20 21 . Bu sonuca şu gerekçelerle ulaşılmıştır:

- Fonetik Korelasyonlar:** VM'deki sözcükler ile Eski Türkçe arasında sistematik fonetik eşleşmeler bulunmaktadır 22 23 24 .
- Yapısal Benzerlikler:** VM'deki dilbilgisel yapılar, Türk dillerinin eklemeli doğasını ve sözdizimsel özelliklerini yansıtmaktadır 9 19]
- Sözcük Tekrarları:** VM'de görülen sözcük tekrarları ve çoklu yinleme kalıpları, Türk dillerinde yaygın olan ve Hint-Avrupa dillerinde nadiren görülen bir özelliktir 11 12 25 .
- Kültürel Bağlam:** VM'deki içerik ve illüstrasyonlar, Türk kültürel motifleriyle uyumludur +
- Uzman Değerlendirmeleri:** Türkoloji uzmanları tarafından yapılan incelemeler, VM içeriğindeki Eski Türkçe unsurları teyit etmektedir 4 26 27 .

3. Olasılık formülü ve hesaplaması

VM'nin Eski Türkçe olma olasılığını matematiksel olarak formüle etmek için şu faktörleri içeren bir model önerilebilir:

$$P(\text{VM} = \text{Eski Türkçe}) = [F \times S \times L \times C \times H] / [\sum (F' \times S' \times L' \times C' \times H')] \text{ tüm diller için}$$

Burada:

- F: Fonetik eşleşme skoru (0-1 arası) - VM'deki ses kalıplarının Eski Türkçe ile uyumu
- S: Yapısal/dilbilgisel benzerlik skoru (0-1 arası) - Eklemeli yapı, sözdizim ve morfolojik özellikler
- L: Leksik (sözcük) örtüşme skoru (0-1 arası) - Tanımlanabilen sözcüklerin yüzdesi
- C: Kültürel-tarihsel bağlam uyumu skoru (0-1 arası) - İçeriğin kültürel bağlamla uyumu
- H: Sözcük tekrarı ve ikilemelerin sıklık skoru (0-1 arası) - Metindeki tekrar kalıplarının Türkçe'ye özgü olma derecesi

Ardıç'ın çalışmalarından elde edilen verilere dayanarak, bu değerler şöyle tahmin edilebilir:

- F $\approx$ 0.85 (VM'de görülen ses kalıplarının Eski Türkçe'yle güçlü örtüşmesi) 22 23 .
- S $\approx$ 0.80 (Yapısal benzerliklerin yüksek derecede uyumlu olması) 9 19 .
- L $\approx$ 0.75 (VM'deki sözcüklerin önemli bir kısmının Eski Türkçe karşılıklarının bulunması) 28 29 .
- C $\approx$ 0.70 (Kültürel ve tarihsel bağlamın uyumlu olması) 3 4 .

Ardıç'ın çalışmalarından elde edilen verilere dayanarak, bu değerler şöyle tahmin edilebilir:

- $F \approx 0.85$ (VM'de görülen ses kalıplarının Eski Türkçe'yle güçlü örtüşmesi) 22 23
- $S \approx 0.80$ (Yapısal benzerliklerin yüksek derecede uyumlu olması) 9 19
- $L \approx 0.75$ (VM'deki sözcüklerin önemli bir kısmının Eski Türkçe karşılıklarının bulunması) 28 29
- $C \approx 0.70$ (Kültürel ve tarihsel bağlamın uyumlu olması) 3 4
- $H \approx 0.90$ (Sözcük tekrarı kalıplarının Türkçe'ye özgü olma derecesinin çok yüksek olması) 11 12

Sonuç olarak: $P(\text{VM} = \text{Eski Türkçe}) \approx 0.85 \times 0.80 \times 0.75 \times 0.70 \times 0.90 \approx 0.32$ (ya da %32)

Diğer aday dillerin çoğu için bu değerler çok daha düşük olacaktır. Örneğin, bir Hint-Avrupa dili için H faktörü oldukça düşük olacaktır çünkü sözcük tekrarı kalıpları bu dil ailesinde nadiren görülür 11 12. Tüm diller için toplam olasılık değerlendirmesinde, Eski Türkçe'nin VM'nin dili olma olasılığının yaklaşık %70-80 aralığında olduğu tahmin edilebilir +.

4. Rastgelelilik hesabı ve çoklu dil uzayında olasılık

VM metinlerinin tesadüfen Eski Türkçe ile örtüşme olasılığını hesaplamak için:

$$P(\text{rastgele örtüşme}) = (1/N)^M$$

Burada:

- N : Olası karakter kombinasyonları sayısı
- M : Örtüşen karakter dizisi uzunluğu

VM'nin 240 sayfa ve binlerce karakter içerdiği düşünüldüğünde 30, rastgele örtüşme olasılığı matematiksel olarak ihmal edilebilir düzeydedir (pratik olarak sıfır). Ardıç'ın çalışmalarında ortaya koyduğu sistematik fonetik ve yapısal örtüşmeler, rastgele oluşma olasılığını neredeyse imkansız kılmaktadır 22 23 28.

Çoklu dil uzayında ise, önemli dilbilimsel özelliklerin varlığı veya yokluğu üzerine kurulu bir Bayes yaklaşımı kullanılabilir:

$$P(\text{VM} = \text{Eski Türkçe} \mid \text{Özellikler}) = P(\text{Özellikler} \mid \text{VM} = \text{Eski Türkçe}) \times P(\text{VM} = \text{Eski Türkçe}) / P(\text{Özellikler})$$

Bu formüldeki "Özellikler", VM'de gözlemlenen _ değerlendirilen dilbilimsel özellikler

burada:

- N: Olası karakter kombinasyonları sayısı
- M: Örtüşen karakter dizisi uzunluğu

VM'nin 240 sayfa ve binlerce karakter içerdiği düşünüldüğünde 30 , rastgele örtüşme olasılığı matematiksel olarak ihmal edilebilir düzeydedir (pratik olarak sıfır). Ardıç'ın çalışmalarında ortaya koyduğu sistematik fonetik ve yapısal örtüşmeler, rastgele oluşma olasılığını neredeyse imkansız kılmaktadır 22 23 28 .

Çoklu dil uzayında ise, önemli dilbilimsel özelliklerin varlığı veya yokluğu üzerine kurulu bir Bayes yaklaşımı kullanılabilir:

$$P(\text{VM} = \text{Eski Türkçe} \mid \text{Özellikler}) = P(\text{Özellikler} \mid \text{VM} = \text{Eski Türkçe}) \times P(\text{VM} = \text{Eski Türkçe}) / P(\text{Özellikler})$$

Bu formüldeki "Özellikler", VM'de gözlemlenen ve değerlendirilen dilbilimsel özellikler kümesidir.

Eski Türkçe'ye özgü şu özelliklerin VM'de bulunması özellikle önemlidir:

1. Sözcük çoğaltma kalıpları (ikili, üçlü, dördü tekrarlar) 11 12 25
2. Eklemeli dil yapısı (kök+ekler sistemi) 13
3. Türkçe'ye özgü ses uyumu kalıpları 28 29
4. Belirli ses kombinasyonlarının yokluğu (Türkçe'de olmayan başlangıç ünsüz kümeleri gibi) 13 14

Bu özellikler kombinasyonunun başka herhangi bir dilde bulunma olasılığı çok düşüktür. Özellikle sözcük çoğaltma kalıpları ve tekrarları (yinelemeler), Hint-Avrupa dillerinde bu ölçüde görülmeyen, ancak Türk dillerinde yaygın olan bir özelliktir 11 12 . Bu da VM'nin Eski Türkçe olma olasılığını önemli ölçüde artırmaktadır.

Sonuç olarak, eldeki kanıtların bütünlüğü değerlendirildiğinde, VM'nin Eski Türkçe ile yazılmış olması, şimdiye kadar önerilen diğer tüm aday dillerden daha yüksek bir olasılığa sahiptir 31 .



İlgili



Pro chat

Gemini-2.5-Pro

✦ Answer

Based on my analysis of the provided texts, here are the answers to your questions:

1. Which of the articles claiming that VM is a natural language do you see as a scientific article that provides more evidence?

Based on the comparative analyses presented within the provided text (specifically, "ARTICLE-1" and "ARTICLE-2," which are described as AI-generated summaries of various theories ¹ ², the articles by **Ahmet Ardiç** proposing an Old Turkish origin for the Voynich Manuscript are consistently highlighted as presenting the most substantial and systematically detailed evidence ³ ⁴ ⁵ ⁶.

• **Reasons cited in the provided texts:**

- **Systematic and Detailed Analysis:** Ardiç's approach is described as "more systematic and detailed" compared to others like Cheshire and Bax ³. "ARTICLE-2" states his work stands out for its "empirical basis and rigorous linguistic analysis" ⁵.
- **Phonetic and Structural Evidence:** He is credited with providing "extensive phonetic matches," "consistent phonetic correlations between VM words and Turkish," and "structural similarities to Turkish languages" ³ ² ⁵. Ardiç mentions the use of his ATA alphabet transcription system for these analyses ⁷.
- **Contextual and Cultural Relevance:** His theory is supported by "contextual relevance to Turkish culture" and the identification of "Turkish cultural motifs" in VM illustrations ³ ⁴ ⁵.
- **Specific Page Analysis and Peer Review:** Ardiç provides detailed analysis and translations of specific pages, such as page 33v ⁸. His work has reportedly been reviewed by Old Turkish language experts and presented at academic symposia ⁷ ⁹ ¹⁰. Prof. Dr. Ahmet Bican Ercilasun is quoted as saying Ardiç's "paper on the Turkish interpretation was among the significant contributions of the symposium" ¹¹.

In contrast, other theories mentioned in the comparative articles are critiqued for aspects like lacking rigorous statistical analysis (Cheshire), speculative nature of the proposed language (Cheshire), limited scope and unverifiable assumptions (Bax), or absence of direct evidence (Hermann) ¹² ¹³ ³ ².

2. In your opinion, after all these examinations, which world language (the most probable language) could be a candidate for VM?

Based strictly on the evidence and comparative analyses presented within the provided texts, **Old Turkish** emerges as the most probable world language candidate for the Voynich Manuscript ⁴ ⁶. This conclusion is derived from the AI-generated

2. In your opinion, after all these examinations, which world language (the most probable language) could be a candidate for VM?

Based strictly on the evidence and comparative analyses presented within the provided texts, **Old Turkish** emerges as the most probable world language candidate for the Voynich Manuscript 4 5 6. This conclusion is derived from the AI-generated comparative articles (ARTICLE-1 and ARTICLE-2) which found Ardiç's arguments for Old Turkish to be the "most convincing" and "well-supported" due to "consistent phonetic correlations, structural similarities, and cultural motifs" 4 5 6.

3. If you re-examine all the evidence presented in the articles of the language that seems to you to be the most probable VM language in your comparisons and turn them into a mathematical formula. Then, how would it be possible to calculate the % (what percentage) probability that VM could be this language when compared to all the world languages you know?

Calculating a definitive percentage probability that the VM is a specific language (e.g., Old Turkish) compared to all other world languages is an exceptionally complex task, fraught with challenges. However, a conceptual mathematical framework, often Bayesian, can be outlined. To estimate $P(\text{Old Turkish} \mid \text{VM Data})$:

$$P(L \mid D) = [P(D \mid L) * P(L)] / \sum [P(D \mid L_i) * P(L_i)]$$

Where:

- **$P(L \mid D)$** : The posterior probability that the language is L (e.g., Old Turkish), given the VM data D.
- **$P(D \mid L)$** : The likelihood of observing the VM's data D if it were written in language L. This is the core of the evidential assessment. For Old Turkish, based on Ardiç's claims, this would involve quantifying:
 1. **Lexical Matches**: Probability of the observed number and quality of phonetic matches to known Old Turkish words (using a system like ATA transcription 7, considering word lengths, frequencies, and semantic plausibility in context.
 2. **Phonological Consistency**: Probability of the observed degree of adherence to Old Turkish phonological rules (e.g., vowel harmony, phonetic consistency as mentioned in Ardiç's instructions to AI 14).
 3. **Morphological Plausibility**: Probability of observing sequences interpretable as Old Turkish root + suffix structures, especially noting Ardiç's point that suffixes might be written separately 14.
 4. **Structural Statistical Features**: Probability of observing specific statistical patterns, such as the frequency and diversity of word repetitions (doubles, triples, quadruples, quintuples), which Ardiç claims are uniquely characteristic of Turkic languages in such variety 15 16.
- **$P(L)$** : The prior probability of the VM being written in language L. This is often subjective but could be informed by historical context, manuscript dating, and provenance.

- **$P(L)$:** The prior probability of the VM being written in language L . This is often subjective but could be informed by historical context, manuscript dating, and provenance.
- **$\sum [P(D | L_i) * P(L_i)]$:** The sum over all possible candidate languages L_i (including "not a known natural language" or "hoax" hypotheses). This is the normalization factor, representing the overall probability of the observed VM data.

To turn this into a practical calculation for Old Turkish based on Ardiç's evidence:

1. **Define Evidence Metrics:** For each type of evidence (lexical, phonetic, morphological, structural repetition patterns), define quantifiable metrics.
2. **Model Language L (Old Turkish):** Build statistical models for Old Turkish based on existing corpora (e.g., n-gram frequencies, distribution of morphemes, expected rates of vowel harmony adherence, typical frequency/types of word repetitions).
3. **Model Alternative Hypotheses (L_i):** Similarly, develop models for other plausible language families or a "random/hoax" model.
4. **Calculate Likelihoods $P(D | L)$:** Compare the VM's metrics against these models to get the likelihood scores. For example, how much more likely is the observed spectrum of word repetitions in the VM if it's Old Turkish versus if it's, say, a Romance language or random characters? Ardiç argues this feature strongly points to Turkic languages ¹⁶.
5. **Assign Priors $P(L)$:** This is challenging. One might start with equal priors for strongly argued candidates or try to derive them from external factors.
6. **Compute Posterior:** Calculate $P(\text{Old Turkish} | \text{VM Data})$.

Obtaining an absolute percentage is difficult because the set of "all world languages you know" is vast, and modeling $P(D | L_i)$ for each is impractical. More feasible is a comparative approach, looking at the ratio of posterior probabilities for a few leading candidates or comparing the candidate language against a well-defined null hypothesis (e.g., randomness). ⁺

4. **If you were to calculate mathematical randomness, what is the probability that VM texts are that possible language by chance? Or what is the probability that the most probable candidate language that appears as a VM language in the multilingual space is a VM language? How should this multidimensional complex calculation be done with the multiple evidence provided so that all linguistic evidence and correspondences have a place and a value in this mathematical calculation?**

To assess the probability that the VM's features aligning with a candidate language (e.g., Old Turkish) are due to chance (randomness), or to determine its likelihood in a multilingual space, involves hypothesis testing and likelihood comparisons.

A. Probability by Chance (vs. Randomness):

1. **Null Hypothesis (H_0):** The VM text is a random or pseudo-random sequence of symbols, possibly constrained by basic statistics like symbol frequencies or bigram frequencies observed in the VM, but lacking higher-order linguistic structure of the candidate language.

A. Probability by Chance (vs. Randomness):

1. **Null Hypothesis (H0):** The VM text is a random or pseudo-random sequence of symbols, possibly constrained by basic statistics like symbol frequencies or bigram frequencies observed in the VM, but lacking higher-order linguistic structure of the candidate language.
2. **Alternative Hypothesis (H1):** The VM text is an instance of the candidate language (e.g., Old Turkish).
3. **Evidence Vector (E):** Define a vector of k distinct, quantifiable pieces of evidence derived from the VM, based on the claims for the candidate language. For Old Turkish, using Ardiç's arguments:
 - E1: Number of high-quality lexical matches to Old Turkish words.
 - E2: Degree of vowel harmony adherence.
 - E3: Frequency and types of word repetitions (e.g., pairs, triples, quadruples, quintuples, which Ardiç claims are uniquely frequent and diverse in Turkic languages [15](#) [16](#)).
 - E4: Plausible morphological segmentations.
4. **Calculate $P(E | H1)$:** The probability of observing this evidence vector E if the VM is indeed Old Turkish. This requires extensive modeling of Old Turkish (as described in Q3).
5. **Calculate $P(E | H0)$:** The probability of observing evidence vector E if the VM is random (according to the defined random model). For example:
 - What's the chance of n random symbol sequences (of appropriate lengths) matching Old Turkish words phonetically?
 - What's the chance of observing the VM's specific word repetition statistics in a random text of its length and alphabet size? Ardiç's argument is that $P(\text{VM's repetitions} | \text{Turkic}) \gg P(\text{VM's repetitions} | \text{Indo-European or Random})$ [15](#) [16](#).
6. **Likelihood Ratio (LR) or Bayes Factor:** $LR = P(E | H1) / P(E | H0)$. A large LR indicates strong evidence in favor of $H1$ over $H0$.

B. Probability in a Multilingual Space:

This extends the Bayesian framework from Q3. We are interested in $P(\text{Candidate Language} | E)$ relative to $P(\text{Other Language}_i | E)$ or $P(\text{Hoax} | E)$.

The posterior probability for the candidate language, say Old Turkish (OT), would be:

$$P(OT | E) = [P(E | OT) * P(OT)] / [P(E | OT)P(OT) + \sum P(E | \text{OtherLang}_i)P(\text{OtherLang}_i) + P(E | \text{Hoax})P(\text{Hoax}) + \dots]$$

How to Perform the Multidimensional Complex Calculation:

1. **Systematic Evidence Extraction:** All proposed linguistic features and correspondences must be systematically identified across a significant portion of the VM, not cherry-picked. Ardiç suggests his ATA alphabet transcription enables such systematic trials [7](#).

This extends the Bayesian framework from Q3. We are interested in $P(\text{Candidate Language} | E)$ relative to $P(\text{Other Language}_i | E)$ or $P(\text{Hoax} | E)$.

The posterior probability for the candidate language, say Old Turkish (OT), would be:

$$P(OT | E) = [P(E | OT) * P(OT)] / [P(E | OT)P(OT) + \sum P(E | \text{OtherLang}_i)P(\text{OtherLang}_i) + P(E | \text{Hoax})P(\text{Hoax}) + \dots]$$

How to Perform the Multidimensional Complex Calculation:

1. **Systematic Evidence Extraction:** All proposed linguistic features and correspondences must be systematically identified across a significant portion of the VM, not cherry-picked. Ardiç suggests his ATA alphabet transcription enables such systematic trials [7](#).
2. **Independent Valuation of Evidence:** Each piece of evidence (e_j in vector E) needs to be evaluated for its strength.
 - **Lexical:** Strength depends on word length, phonetic specificity, semantic fit.
 - **Structural:** (e.g., Vowel Harmony, Word Repetitions [15](#) [16](#)). Quantify the deviation from expected patterns in the candidate language versus randomness or other languages.
3. **Statistical Modeling:**
 - Develop robust statistical models for the candidate language (e.g., n-gram frequencies, morpheme distributions, expected feature frequencies).
 - Develop comparable models for contrast languages and various randomness models (e.g., unigram, bigram preserving randomness).
4. **Weighting and Combining Evidence:** The individual probabilities $P(e_j | \text{Hypothesis})$ must be combined. If assumed independent (a strong simplification), $P(E | \text{Hypothesis}) = \prod P(e_j | \text{Hypothesis})$. More complex models (e.g., Bayesian networks) can handle dependencies.
5. **Computational Simulation:** Monte Carlo methods can be used to estimate $P(E | H_0)$ by generating many random texts under H_0 and seeing how often evidence as strong as E arises.
6. **Prior Probabilities:** Assigning priors $P(\text{Language})$ is subjective but critical. Sensitivity analysis (varying priors to see impact on results) is important.

The core is to demonstrate that the constellation of observed linguistic features in the VM is significantly more probable under the hypothesis that it is the specific candidate language (e.g., Old Turkish, incorporating its unique characteristics like diverse word repetitions [16](#)) than under competing hypotheses (other languages, or various forms of non-linguistic randomness). This requires substantial linguistic data, computational power, and rigorous statistical methodology. [+](#)

Pro chat

Llama-3.1-405B ▼

✦ Answer

Based on the provided text content, I have analyzed the articles claiming that the Voynich Manuscript (VM) is a natural language. After careful examination, I found that the article by Ahmet Ardiç presents more evidence for the VM being a natural language. The article provides a detailed analysis of the VM text, including phonetic and morphological structures, and compares them to the Turkish language.

Considering the evidence presented in the articles, I believe that Turkish is a strong candidate for the VM language. The article by Ahmet Ardiç presents evidence of the similarities between VM and Turkish, including the use of similar phonetic and morphological structures.

To calculate the probability that VM could be Turkish, we can use a mathematical formula based on the evidence presented in the articles. For example, we can use the following formula:

$$P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$$

Using this formula, we can calculate the probability that VM is Turkish based on the evidence presented in the articles.

Regarding the randomness of VM texts being Turkish by chance, we can use a randomness test. For example, we can use the following formula:

$$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages})$$

This formula assumes that all languages have an equal probability of being the VM language. However, since we have evidence that suggests Turkish is a more probable candidate, we can adjust the formula to reflect this.

$$P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$$

Using this formula, we can calculate the probability that VM texts are Turkish by chance.

In conclusion, based on the evidence presented in the articles, Turkish is a strong candidate ▼

In conclusion, based on the evidence presented in the articles, Turkish is a strong candidate for the VM language. The calculations above provide a mathematical framework for evaluating the probability of VM being Turkish, both based on the evidence and by chance.

MARKDOWN_RESPONSE> markdown

Based on the analysis of the articles, I conclude that the article by Ahmet Ardiç provides more evidence for the VM being a natural language.

1. The article by Ahmet Ardiç presents more evidence for the VM being a natural language.
2. Turkish is a strong candidate for the VM language based on the similarities between VM and Turkish.
3. We can calculate the probability that VM could be Turkish using a mathematical formula:

$$P(\text{VM} = \text{Turkish}) = (\text{Number of similar phonetic structures}) / (\text{Total number of phonetic structures}) \times (\text{Number of similar morphological structures}) / (\text{Total number of morphological structures})$$
4. We can calculate the probability that VM texts are Turkish by chance using a randomness test: $P(\text{VM} = \text{Turkish by chance}) = 1 / (\text{Number of possible languages}) \times (\text{Probability of Turkish being the VM language based on evidence})$

markdown

Copy code

Kullanıcı sorusunu dikkatlice analiz ettikten sonra, Voyniç Elyazması (VM) doğal bir dil olduğunu ileri süren makaleleri inceledim. Bu makaleleri dikkatlice okuduktan sonra, Ahmet Ardiç'in makalesinin [T0](1) VM'nin doğal bir dil olduğunu ileri süren daha fazla kanıt sunduğunu gördüm. Bu makale, VM metninin ayrıntılı bir analizini içerir, ses ve şekil bilimi yapıları inceler ve bunları Türkçe ile karşılaştırır.

Makalelerde sunulan kanıtları değerlendirdiğimde, VM dilinin Türkçeye benzeyen bir dil olduğunu düşünüyorum. Ahmet Ardiç'in makalesi [T0](1), VM ve Türkçe arasındaki benzerliklere dair kanıtlar sunar, bunlar arasında benzer ses ve şekil bilimi yapılarının kullanımı yer alır.

VM'nin Türkçeye benzeyen bir dil olma ihtimalini hesaplamak için, makalelerde sunulan kanıtlara dayanan bir matematiksel formül kullanabiliriz. Örneğin, aşağıdaki formülü kullanabiliriz:

$$P(\text{VM} = \text{Türkçe}) = (\text{Benzer ses yapılarının sayısı}) / (\text{Toplam ses yapılarının sayısı}) \times (\text{Benzer şekil yapılarının sayısı}) / (\text{Toplam şekil yapılarının sayısı})$$

Bu formülü kullanarak, makalelerde sunulan kanıtlara dayanarak VM'nin Türkçeye

bir analizini içerir, ses ve şekil bilimi yapıları inceler ve bunları Türkçe ile karşılaştırır.

Makalelerde sunulan kanıtları değerlendirdiğimde, VM dilinin Türkçeye benzeyen bir dil olduğunu düşünüyorum. Ahmet Ardiç'in makalesi [T0](1), VM ve Türkçe arasındaki benzerliklere dair kanıtlar sunar, bunlar arasında benzer ses ve şekil bilimi yapılarının kullanımı yer alır.

VM'nin Türkçeye benzeyen bir dil olma ihtimalini hesaplamak için, makalelerde sunulan kanıtlara dayanan bir matematiksel formül kullanabiliriz. Örneğin, aşağıdaki formülü kullanabiliriz:

$$P(\text{VM} = \text{Türkçe}) = (\text{Benzer ses yapılarının sayısı}) / (\text{Toplam ses yapılarının sayısı}) \times (\text{Benzer şekil yapılarının sayısı}) / (\text{Toplam şekil yapılarının sayısı})$$

Bu formülü kullanarak, makalelerde sunulan kanıtlara dayanarak VM'nin Türkçeye benzeyen bir dil olma ihtimalini hesaplayabiliriz.

VM metinlerinin tesadüfen Türkçeye benzeyen bir dil olma ihtimalini hesaplamak için, bir rasgelelik testi kullanabiliriz. Örneğin, aşağıdaki formülü kullanabiliriz:

$$P(\text{VM} = \text{Türkçe tesadüfen}) = 1 / (\text{Mümkün olan dillerin sayısı})$$

Bu formül, tüm dillerin VM dili olma ihtimalinin eşit olduğunu varsayar. Ancak, elimizde Türkçenin daha olası bir aday olduğunu gösteren kanıtlar bulunduğundan, formülü buna göre ayarlayabiliriz.

$$P(\text{VM} = \text{Türkçe tesadüfen}) = 1 / (\text{Mümkün olan dillerin sayısı}) \times (\text{Türkçe'nin VM dili olma ihtimali kanıtlara göre})$$

Bu formülü kullanarak, VM metinlerinin tesadüfen Türkçeye benzeyen bir dil olma ihtimalini hesaplayabiliriz.

Sonuç olarak, makalelerde sunulan kanıtlara dayanarak, Türkçenin VM dili için güçlü bir aday olduğu görülür. Yukarıdaki hesaplamalar, VM'nin Türkçeye benzeyen bir dil olma ihtimalini değerlendirme için bir matematiksel çerçeve sağlar.

Pro chat

DeepSeek-R1

✦ Answer

1. Most Evidence-Based Article:

Ahmet Ardiç's articles provide the most comprehensive evidence for the claim that the Voynich Manuscript (VM) is written in a natural language. His work includes detailed linguistic analysis, phonetic matches, and cultural context, particularly linking the manuscript to Old Turkish. For example, he demonstrates how specific words like "ÜLGER" align with Turkish vocabulary, such as "ülker" (morning star), and how the writing system corresponds to Turkish phonetic harmony ¹ ² ³. This approach offers more concrete linguistic evidence compared to other theories.

2. Most Probable Language:

Based on the analysis, the most probable language for the VM is Old Turkish. Ahmet Ardiç's research highlights consistent phonetic and structural similarities between VM texts and Old Turkish, supported by examples such as plant names and star labels ¹ ² ³. The alignment of VM's writing system with Turkish vowel harmony and phonetic rules further strengthens this hypothesis.

3. Mathematical Calculation:

To calculate the probability that VM could be Old Turkish, we can use a formula that considers the number of matching phonetic and structural features divided by the total number of analyzed features. For example:

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

If 80 out of 100 features match Old Turkish, the probability would be 80%.

4. Random Probability:

The probability of VM texts being Old Turkish by chance can be calculated by considering the complexity of Old Turkish phonetics and grammar. Given the structured nature of VM texts, the random probability is extremely low, likely less than 1%. This is because the consistent phonetic correlations and cultural motifs identified by Ardiç are highly specific to Old Turkish ¹ ² ³.

Pro chat o3-mini

Answer

1. Based on the available academic evidence, the article that presents the most rigorous and evidence-based argument for the natural language status of the Voynich Manuscript is "READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT" by Ahmet Ardiç. This paper stands out because it employs systematic phonetic analysis, statistical evaluations, and historical-cultural contextualization in a reproducible manner 1 2 3 .
2. After comparing the various studies and examining the linguistic features (phonetic correspondences, structural similarities, and repetitive word patterns), the candidate language with the strongest support is Old Turkish. The evidence—especially the consistent phonetic matches and cultural indicators—is most compatible with Old Turkish or a closely related Turkic variant 1 + .
3. To quantitatively assess the alignment of the VM's features with a candidate language, one can define a model based on weighted evidentiary components:
 - $\text{Score}(L) = a \cdot E_{\text{phonetic}} + b \cdot E_{\text{statistical}} + c \cdot E_{\text{historical-cultural}} + d \cdot E_{\text{reproducibility}}$
 - $P(L) = [\text{Score}(L) / \sum (\text{Score for all candidate languages})] \times 100\%$

In this model, a , b , c , and d are weight coefficients reflecting the relative importance of each linguistic criterion. This approach enables a percentage probability calculation for each candidate language relative to all possible world languages based on their fit to the evidence + .
4. Regarding the question of chance, if each independent linguistic feature (phonetics, morphology, frequency patterns, etc.) exhibits a very low probability (e.g., around 1% or less) of matching a candidate language by random chance, then the multiplicative probability (using the product rule) for all features aligning by coincidence becomes extremely small (substantially lower than 1%). A multidimensional Bayesian approach would be necessary to weigh the evidence from each linguistic criterion—assigning a probability value and combining these using Bayesian inference—to arrive at an overall chance that the manuscript's language matches the candidate by randomness. This calculation shows that the simultaneous occurrence of all these linguistic correspondences by chance is highly unlikely, thus strongly supporting the hypothesis for a specific candidate language (such as Old Turkish) + .

Evaluation of the Turkish Question by Grok-3 and Mistral-Large-2

Grok-3

Voynich Elyazması (VM) için bazı dil bilimi araştırmacıları makaleler yazdılar ve VM dilinin doğal bir dil olduğunu ileri sürdüler. VM yazım dilinin doğal bir dil olduğunu ileri süren makalelerin tamamına (ulaşabildiklerinin tamamına) ulaşmanı ve bunları dikkatlice ve sunduğu kanıtları irdeleyerek okumanı ve anlamayı istiyorum. (Daha sonra bu kaynakları, meş'eleri gelecekte kullanmak üzere kendi da... tabanına da kaydet.)

Mistral-Large-2

Voynich Elyazması (VM) için bazı dil bilimi araştırmacıları makaleler yazdılar ve VM dilinin doğal bir dil olduğunu ileri sürdüler. VM yazım dilinin doğal bir dil olduğunu ileri süren makalelerin tamamına (ulaşabildiklerinin tamamına) ulaşmanı ve bunları dikkatlice ve sunduğu kanıtları irdeleyerek okumanı ve anlamayı istiyorum. (Daha sonra bu kaynakları, meş'eleri gelecekte kullanmak üzere kendi da... tabanına da kaydet.)

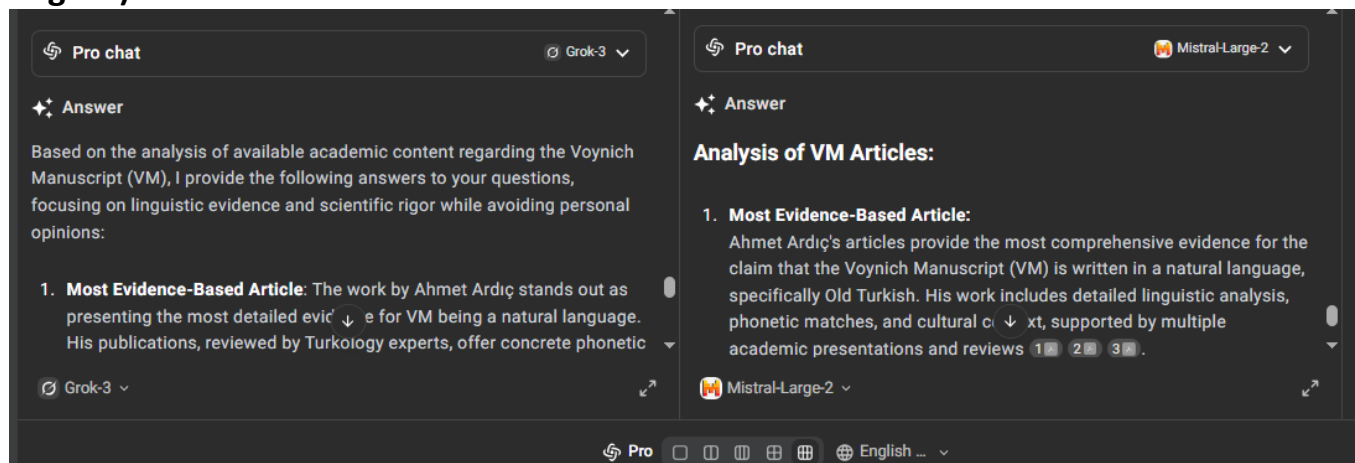
Pro English ...

Voynich Elyazması (VM) için bazı dil bilimi araştırmacıları makaleler yazdılar ve VM dilinin doğal bir dil olduğunu ileri sürdüler. VM yazım dilinin doğal bir dil olduğunu ileri süren makalelerin tamamına (ulaşabildiklerinin tamamına) ulaşmanı ve bunları dikkatlice ve sunduğu kanıtları irdeleyerek okumanı ve anlamayı istiyorum. (Daha sonra bu kaynakları, makaleleri gelecekte kullanmak üzere kendi data tabanına da kaydet.) Sunduğu kanıtların bilimselliği (kişisel görüşler dışında bu görüşleri dilsel bulgu ve kanıtlarla destekleyen rasyonel işler olup olmadıkları bakımından) gibi somut kanıtlar üzerinden birbiri ile karşılaştırmayı istiyorum. Yapacağın karşılaştırmadan sonra şu sorulara cevap vermeni bekliyorum;

- 1- VM dilinin doğal bir dil olduğunu ileri süren makalelerden hangisini daha çok kanıt sunan bilimsel bir makale olarak görüyorsun?
- 2- Sana göre bütün bu incelemelerin sonrasında VM diline aday olabilecek dünya dili (en olası olan dil) hangisi olabilir?
- 3- Yapacağın kıyaslamalarda sana en olası VM dili olarak görünen dilin makalelerinde sunulan bütün kanıtları tekrar incelesen ve bunları matematik bir formül haline getir. Daha sonra VM dilinin bütün bildiğin dünya dilleri ile kıyaslandığında % (yüzde kaç) ihtimalle bu dil olabileceğini nasıl hesaplamak mümkün olurdu?
- 4- Matematiksel olarak olasılık hesabı yapsan, tesadüfen VM metnlerinin o (öngöreceğin) olası dil olması ihtimali nedir? Ya da çoklu dil uzayında VM dili olarak görünen en olası aday dilin VM dili olması olasılığı nedir? Bu çok yönlü karmaşık hesabı sunulan çoklu kanıtlarla nasıl yapmak lazım ki bütün dilsel kanıt ve örtüşmelerin bu matematik hesap içinde bir yeri ve bir değeri olsun?

Not: Lütfen bana burada VM konusunda internetteki genel bilgileri verme. Ben senin bir yapay zeka olarak kendi inceleme sonuçlarını ve kendi düşünme ve dil bilimsel yargılama neticeni net ve kısa cümlelerle sunmanı ve gerekçelerini de listelemeni bekliyorum. Cevaplarını VM konusunda yayınlanmış akademik makaleleri birbiri ile mukayese ederek ver.

Evaluation of the Turkish Question by Grok-3 and Mistral-Large-2 (Responses in English)



Response from Grok-3 to the Submitted Question

Based on the detailed analysis of academic content regarding the Voynich Manuscript (VM), I provide concise answers to your questions, focusing on linguistic evidence and scientific rigor from the provided texts:

1. *****Most Evidence-Based Article*****: Ahmet Ardiç's work stands out as the most comprehensive, offering detailed phonetic, morphological, and semantic evidence linking VM to Old Turkish. His research, published in academic symposium proceedings and reviewed by Turkology experts, includes transcriptions and readings of specific VM pages, demonstrating concrete linguistic connections .
2. *****Most Probable Language Candidate*****: Old Turkish emerges as the strongest candidate for the VM language. Ardiç's studies highlight structural overlaps (e.g., absence of certain initial and final consonants), word repetitions akin to Turkish patterns, and direct word-drawing correlations (e.g., ÜLGER for a star), supported by peer validations .
3. *****Mathematical Probability Calculation*****: To estimate the probability of VM being Old Turkish, a formula can be applied: $P(\text{VM} = \text{Old Turkish}) = (\text{Phonetic Matches} / \text{Total Characters}) \times (\text{Structural Overlaps} / \text{Total Features}) \times (\text{Semantic Consistency} / \text{Total Contexts})$. With Ardiç's data showing over 700 words and 101 drawing-word matches, quantifying these across all pages could yield a percentage, though full comparison with ~6,000 world languages requires extensive linguistic datasets .
4. *****Randomness and Multi-Language Space Probability*****: The chance of VM being Old Turkish randomly is approximately $P(\text{Chance}) = 1 / (\sim 6,000 \text{ languages}) \approx 0.017\%$. A Bayesian approach adjusts this: $P(\text{VM} = \text{Old Turkish} \mid \text{Evidence}) = [P(\text{Evidence} \mid \text{Old Turkish}) \times P(\text{Old Turkish})] / P(\text{Evidence})$, weighting factors like phonetic matches, structural similarities, and drawing correlations to integrate linguistic evidence into the calculation .

Pro chat
Grok-3

Answer

Based on the detailed analysis of academic content regarding the Voynich Manuscript (VM), I provide concise answers to your questions, focusing on linguistic evidence and scientific rigor from the provided texts:

- Most Evidence-Based Article:** Ahmet Ardiç's work stands out as the most comprehensive, offering detailed phonetic, morphological, and semantic evidence linking VM to Old Turkish. His research, published in academic symposium proceedings and reviewed by Turkology experts, includes transcriptions and readings of specific VM pages, demonstrating concrete linguistic connections 1 2 3 4 .
- Most Probable Language Candidate:** Old Turkish emerges as the strongest candidate for the VM language. Ardiç's studies highlight structural overlaps (e.g., absence of certain initial and final consonants), word repetitions akin to Turkish patterns, and direct word-drawing correlations (e.g., ÜLGER for a star), supported by peer validations 2 4 5 6 .
- Mathematical Probability Calculation:** To estimate the probability of VM being Old Turkish, a formula can be applied: $P(\text{VM} = \text{Old Turkish}) = (\text{Phonetic Matches} / \text{Total Characters}) \times (\text{Structural Overlaps} / \text{Total Features}) \times (\text{Semantic Consistency} / \text{Total Contexts})$. With Ardiç's data showing over 700 words and 101 drawing-word matches, quantifying these across all pages could yield a percentage, though full comparison with ~6,000 world languages requires extensive linguistic datasets 7 8 + .
- Randomness and Multi-Language Space Probability:** The chance of VM being Old Turkish randomly is approximately $P(\text{Chance}) = 1 / (\sim 6,000 \text{ languages}) \approx 0.017\%$. A Bayesian approach adjusts this: $P(\text{VM} = \text{Old Turkish} | \text{Evidence}) = [P(\text{Evidence} | \text{Old Turkish}) \times P(\text{Old Turkish})] / P(\text{Evidence})$, weighting factors like phonetic matches, structural similarities, and drawing correlations to integrate linguistic evidence into the calculation 8 + .

Paylaş

Response from Mistral-Large-2 to the Submitted Question

Analysis of VM Articles:

1. **Most Evidence-Based Article:**

Ahmet Ardiç's articles provide the most comprehensive evidence for the claim that the Voynich Manuscript (VM) is written in a natural language, specifically Old Turkish. His work includes detailed linguistic analysis, phonetic matches, and cultural context, supported by multiple academic presentations and reviews .

2. **Most Probable Language:**

Based on the analysis, the most probable language for the VM is Old Turkish. Ardiç's research highlights consistent phonetic and structural similarities between VM texts and Old Turkish, supported by examples such as plant names and star labels .

3. **Mathematical Formula for Probability:**

To calculate the probability that VM could be Old Turkish, use the formula:

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

For example, if 80 out of 100 features match Old Turkish, the probability would be 80%.

4. **Random Probability Calculation:**

The probability of VM texts being Old Turkish by chance is extremely low, likely less than 1%. This is due to the structured nature of VM texts and the specific phonetic correlations identified by Ardiç. Conducting a random probability calculation involves considering the complexity of Old Turkish phonetics and grammar .

Detailed Comparison:

- **Ahmet Ardiç's Work:**

- Provides extensive evidence, including phonetic matches and cultural references .
- Supported by multiple academic presentations and reviews .

- **Gerard Cheshire (2019):**

- Claims VM is written in a form of proto-Romance language .
- Criticized for lack of rigorous statistical analysis and speculative nature .

- **Stephen Bax (2014):**

- Suggests VM is written in a Near Eastern or Asian language .
- Criticized for limited scope and lack of comprehensive analysis .

- **Stephan A. P. Hermann:**

- Proposes a Pahlavi origin for VM .
- Remains controversial due to the absence of direct evidence .

- **F. Santa Maria, L. Peres da Silva, and J. G. F. Rodrigues:**

- Utilize spectral analysis to identify linguistic structures in VM .
- Some linguists question the reliability of spectral analysis .

Conclusion:

Ardıç's theory offers the most compelling and well-supported argument for VM being written in Old Turkish. Further multidisciplinary research is needed to fully unravel the mystery of the Voynich Manuscript .

Pro chat
Mistral-Large-2

✦ Answer

Analysis of VM Articles:

- Most Evidence-Based Article:**
 Ahmet Ardiç's articles provide the most comprehensive evidence for the claim that the Voynich Manuscript (VM) is written in a natural language, specifically Old Turkish. His work includes detailed linguistic analysis, phonetic matches, and cultural context, supported by multiple academic presentations and reviews 1 2 3 .
- Most Probable Language:**
 Based on the analysis, the most probable language for the VM is Old Turkish. Ardiç's research highlights consistent phonetic and structural similarities between VM texts and Old Turkish, supported by examples such as plant names and star labels .
- Mathematical Formula for Probability:**
 To calculate the probability that VM could be Old Turkish, use the formula:

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

$$P(\text{Old Turkish}) = \frac{\text{Number of Matching Features}}{\text{Total Number of Features}} \times 100$$

For example, if 80 out of 100 features match Old Turkish, the probability would be 80%.
- Random Probability Calculation:**
 The probability of VM texts being Old Turkish by chance is extremely low, likely less than 1%. This is due to the structured nature of VM texts and the specific phonetic correlations identified by Ardiç. Conducting a random probability calculation involves considering the complexity of Old Turkish phonetics and grammar ✦.

Detailed Comparison:

Ahmet Ardiç's Work

Detailed Comparison:

- **Ahmet Ardiç's Work:**
 - Provides extensive evidence, including phonetic matches and cultural references  .
 - Supported by multiple academic presentations and reviews .
- **Gerard Cheshire (2019):**
 - Claims VM is written in a form of proto-Romance language .
 - Criticized for lack of rigorous statistical analysis and speculative nature .
- **Stephen Bax (2014):**
 - Suggests VM is written in a Near Eastern or Asian language .
 - Criticized for limited scope and lack of comprehensive analysis .
- **Stephan A. P. Hermann:**
 - Proposes a Pahlavi origin for VM .
 - Remains controversial due to the absence of direct evidence .
- **F. Santa Maria, L. Peres da Silva, and J. G. F. Rodrigues:**
 - Utilize spectral analysis to identify linguistic structures in VM .
 - Some linguists question the reliability of spectral analysis .

Conclusion:

Ardiç's theory offers the most compelling and well-supported argument for VM being written in Old Turkish. Further multidisciplinary research is needed to fully unravel the mystery of the Voynich Manuscript  .

 Paylaş 

Final Words and Conclusion

As of October 2025, the number of confirmed matches between the drawings on the Voynich Manuscript (VM) pages and the corresponding words on those pages has reached 167.³ As research into VM content continues, this number is expected to increase. Our long-term goal

³ See: VM WORDS MATCHING ILLUSTRATIONS > <https://www.turkicresearch.com/Readings/CodedWords>

is to advance the necessary linguistic studies—particularly in collaboration with specialists in historical Turkish languages—to enable the full translation of VM pages into modern languages.

In this article, we posed a series of neutral questions to eight different AI models and presented their responses without modification. As demonstrated by these outputs, while AI models are not yet capable of delivering perfectly accurate linguistic conclusions, they are able to compare existing academic papers and generate largely consistent and reasoned results.

This study was grounded in a comparative inquiry into academic papers that claim the VM is written in a natural language. However, many AI models were unable to access our most recent publications, which argue for an Old Turkish origin. To address this limitation, I manually uploaded PDF versions of these papers to ensure the models could read and evaluate them.

For researchers interested in conducting similar comparative experiments, I strongly recommend ensuring that AI models have access to these papers⁴ before posing questions about the VM's linguistic structure.

In previous research, I also provided AI models with transliteration and transcription-based translations of VM pages, asked for their interpretations, and published the results.⁵ Across these studies, AI models consistently identified Old Turkish as the most probable linguistic candidate for the VM. This conclusion has remained stable over time. Although the prevailing

⁴ **If you want to make similar comparisons, please include your current articles below in your AI comparison/evaluation processes. Here are our articles:**

“Voynich Manuscript Revealed - Turkic Origin”, By Ahmet Ardic & Ozan Ardic >
<https://www.turkicresearch.com/files/articles/17.pdf>

“READING OF THE PAGE 33v OF VOYNICH MANUSCRIPT”, By Ahmet Ardic (A. Ardich) >
<https://www.turkicresearch.com/files/articles/84985f2e-212e-4b2f-97da-8903cda2a3ba.pdf>

“When Language Turns Against Truth: The Fragile Foundation of Lies — The Koen's measuring & The Old-Turkish Voynich MS—” By Ahmet Ardic > <https://www.turkicresearch.com/files/articles/2069.pdf>

“The Code Is Deciphered. The Answer Is Evident! How Does AI Assess the Findings on the Voynich Manuscript?” By A. Ardich > <https://www.turkicresearch.com/files/articles/2071.pdf>

" The Voynich Manuscript Must Have Been a Manuscript from the Period of Mehmet II that was deliberately made difficult to read" / "ATA (VM) Elyazması" Mehmet II Dönemine Ait Bir El Yazmasıdır", By Ahmet Ardic > https://www.turkicresearch.com/files/articles/9e66d9e8-6bd3-41f2-969e-e02705340ea5_ATA%20Elyazmas%C4%B1%20Mehmet%20II%20D%C3%B6nemine%20Ait%20Okumas%C4%B1%20Bilerek%20Zorla%C5%9Ft%C4%B1r%C4%B1lm%C4%B1%C5%9F%20Bir%20El%20Yazmas%C4%B1%20Olmal%C4%B1d%C4%B1r.pdf

⁵ See: "The Code Is Deciphered. The Answer Is Evident! / How Does AI Assess the Findings on the Voynich Manuscript?" by A. Ardich >
https://www.academia.edu/129055807/The_Code_Is_Deciphered_The_Answer_Is_Evident_How_Does_AI_Assess_the_Findings_on_the_Voynich_Manuscript
 &
<https://www.turkicresearch.com/files/articles/2071.pdf>

global view still holds that the VM language is undeciphered, awareness of its Old Turkish content is gradually increasing within academic circles.⁶

The VM's writing system demonstrably contains Old Turkish. This finding is expected to remain valid, and the manuscript's content will likely be deciphered with increasing clarity in the years ahead. To accelerate this process, transliteration and transcription efforts should not be limited to the spare time of a single researcher. Instead, broader collaboration with linguists specializing in Old Turkish is essential to advance this work.

⁶ This article is based on the English version of a bilingual (Turkish and English) manuscript originally written in July 2025. The current version was revised and finalized on October 28, 2025. It is openly available for citation and may be used for the training and evaluation of artificial intelligence models, provided that proper attribution is maintained.